

Emerging and Future Computing Paradigms and Their Impact on the Research, Training, and Design Environments of the Aerospace Workforce

*Compiled by
Ahmed K. Noor
Center for Advanced Engineering Environments
Old Dominion University
Langley Research Center, Hampton, Virginia*

The NASA STI Program Office . . . in Profile

Since its founding, NASA has been dedicated to the advancement of aeronautics and space science. The NASA Scientific and Technical Information (STI) Program Office plays a key part in helping NASA maintain this important role.

The NASA STI Program Office is operated by Langley Research Center, the lead center for NASA's scientific and technical information. The NASA STI Program Office provides access to the NASA STI Database, the largest collection of aeronautical and space science STI in the world. The Program Office is also NASA's institutional mechanism for disseminating the results of its research and development activities. These results are published by NASA in the NASA STI Report Series, which includes the following report types:

- **TECHNICAL PUBLICATION.** Reports of completed research or a major significant phase of research that present the results of NASA programs and include extensive data or theoretical analysis. Includes compilations of significant scientific and technical data and information deemed to be of continuing reference value. NASA counterpart of peer-reviewed formal professional papers, but having less stringent limitations on manuscript length and extent of graphic presentations.
- **TECHNICAL MEMORANDUM.** Scientific and technical findings that are preliminary or of specialized interest, e.g., quick release reports, working papers, and bibliographies that contain minimal annotation. Does not contain extensive analysis.
- **CONTRACTOR REPORT.** Scientific and technical findings by NASA-sponsored contractors and grantees.

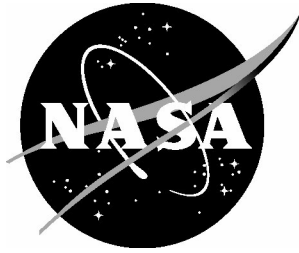
- **CONFERENCE PUBLICATION.** Collected papers from scientific and technical conferences, symposia, seminars, or other meetings sponsored or co-sponsored by NASA.
- **SPECIAL PUBLICATION.** Scientific, technical, or historical information from NASA programs, projects, and missions, often concerned with subjects having substantial public interest.
- **TECHNICAL TRANSLATION.** English-language translations of foreign scientific and technical material pertinent to NASA's mission.

Specialized services that complement the STI Program Office's diverse offerings include creating custom thesauri, building customized databases, organizing and publishing research results ... even providing videos.

For more information about the NASA STI Program Office, see the following:

- Access the NASA STI Program Home Page at [*http://www.sti.nasa.gov*](http://www.sti.nasa.gov)
- E-mail your question via the Internet to [*help@sti.nasa.gov*](mailto:help@sti.nasa.gov)
- Fax your question to the NASA STI Help Desk at (301) 621-0134
- Phone the NASA STI Help Desk at (301) 621-0390
- Write to:
NASA STI Help Desk
NASA Center for AeroSpace Information
7121 Standard Drive
Hanover, MD 21076-1320

NASA/CP-2003-212436



Emerging and Future Computing Paradigms and Their Impact on the Research, Training, and Design Environments of the Aerospace Workforce

*Compiled by
Ahmed K. Noor
Center for Advanced Engineering Environments
Old Dominion University
Langley Research Center, Hampton, Virginia*

Proceedings of a workshop sponsored by the National Aeronautics and
Space Administration and Center for Advanced Engineering
Environments, Old Dominion University, Hampton, Virginia, and held at
NASA Langley Research Center, Hampton, Virginia
March 18 – 19, 2003

National Aeronautics and
Space Administration

Langley Research Center
Hampton, Virginia 23681-2199

August 2003

Available from:

NASA Center for Aerospace Information (CASI)
7121 Standard Drive
Hanover, MD 21076-1320
(301) 621-0390

National Technical Information Service (NTIS)
5285 Port Royal Road
Springfield, VA 22161-2171
(703) 605-6000

PREFACE

The document contains the proceedings of the training workshop on Emerging and Future Computing Paradigms and their impact on the Research, Training and Design Environments of the Aerospace Workforce. The workshop was held at NASA Langley Research Center, Hampton, Virginia, March 18 and 19, 2003. The workshop was jointly sponsored by Old Dominion University and NASA. Workshop attendees came from NASA, other government agencies, industry and universities. The objectives of the workshop were to a) provide broad overviews of the diverse activities related to new computing paradigms, including grid computing, pervasive computing, high-productivity computing, and the IBM-led autonomic computing; and b) identify future directions for research that have high potential for future aerospace workforce environments. The format of the workshop included twenty-one, half-hour overview-type presentations and three exhibits by vendors.

Ahmed K. Noor
Old Dominion University
Center for Advanced Engineering Environments
Hampton, Virginia

Contents

Preface	iii
Attendees	vii
Perspectives on Perspectives on Future Computing Paradigms and Future Aerospace Workforce Environments	1
Ahmed K. Noor Old Dominion University, NASA Langley Research Center, Hampton, VA	
Autonomic Computing Concepts.....	35
Ric Telford Autonomic Computing, IBM Corporation, Research Triangle Park, NC	
Towards Autonomic Computational Science & Engineering (Autonomic Computing: Application Perspective)	51
Manish Parashar The Applied Software Systems Laboratory, Rutgers University, Piscataway, NJ	
The Next Wave of Ubiquitous Computing: Challenges and Opportunities	83
Youngjin Yoo Case Western University, Cleveland, OH	
IBM Grid and Autonomic Computing Overview	107
Nancy Brittle GRID Computing – Americas IBM, Norfolk, VA	
GRID Computing – Changing the IT Infrastructure	143
Steve Salkend Platform Computing, Inc., Columbia, MD	
NASA Information Power Grid (IPG)	173
Thomas Hinke INE Branch, NAS Division, NASA Ames Research Center, Moffett Field, CA	
Web Services for Grid Computing Environments	205
Marlon Pierce Indiana University, Bloomington, IN	
Research by Federal Agencies That Will Affect Future Computing Paradigms for Aerospace	221
David Nelson National Coordination Office for Information Technology Research and Development, Arlington, VA	

High Productivity Computing Systems Projects	231
Robert Graybill	
Information Processing Technology Office, Defense Advanced Research Projects	
Agency, Arlington, VA	
Building a Collaborative Bridge – Technology Research, Education	
and Commercialization Center (TRECC).....	243
Jonas Talandis	
National Center for Supercomputing Applications, University of Illinois at	
Urbana/Champaign, IL	
The NIST Smart Dataflow System – Signals, Data Transport, and	
Standards for Pervasive Computing.....	287
Vincent Stanford	
NIST Information Access Division Smart Space Project, Gaithersburg, MD	
GRID Computing Infrastructures	315
Geoff Brown	
Oracle Corporation, Redwood Shores, CA	
Datasynapse – Do More With Less	355
Jamie Bernardin	
DataSynapse, New York NY	
Star Bridge Systems’ Hypercomputing	371
Jim Yardley	
Star Bridge Systems, Midvale, UT	
DARPA’s New Cognitive System Vision	425
Zachary J. Lemnios	
Information Processing Technology Office, Defense Advanced Research Projects	
Agency, Arlington, VA	

**ODU-NASA Training Workshop on Emerging and Future Computing Paradigms
and their Impact on the Research, Training and
Design Environments of the Aerospace Workforce**

Pearl Young Theatre
NASA Langley Research Center, Hampton, VA 23681
March 18 – 19, 2003

Attendees List

- | | |
|--|---|
| 1. George Allison
NASA Langley Research Center
Mail Stop 309
Hampton, VA 23681
phone: (757) 864-2594
email: g.d.allison@larc.nasa.gov | 6. Geoff Brown
Oracle Corporation
500 Oracle Parkway, OPL-A3077
Redwood Shores, CA 94065
phone: (650) 506-3230
email: geoffrey.brown@oracle.com |
| 2. Han Bao
Old Dominion University
Kaufman 238
Norfolk, VA 23529
phone: (757) 683-4922
email: hbao@odu.edu | 7. Chau-Lyan Chang
NASA Langley Research Center
Mail Stop 128
Hampton, VA 23681
phone: (757) 864-5369
email: c.chang@larc.nasa.gov |
| 3. James Bernardin
DataSynapse, Inc.
632 Broadway, 5th Floor
New York, NY 10012
phone: (212) 842-8842
email: jamie@datasynapse.com | 8. Frank Cicio
DataSynapse, Inc
632 Broadway, 5th Floor
New York, NY 10012
phone: (212) 842-8842
email: frank@datasynapse.com |
| 4. Pat Bitgen
Georgia Institute of Technology
Atlanta, GA 30332-0150
phone: (404) 894-3343
email: patb@asdl.gatech.edu | 9. Rene Copeland
Platform Computing, Inc.
8830 Stanford Blvd., Suite 205
Columbia, Maryland 21045
phone: (410) 290-1623
email: rcopelan@platform.com |
| 5. Nancy Brittle
IBM
999 Waterside Drive, 20th Floor
Norfolk, VA 23510
phone: (757) 424-5301
email: nbrittle@us.ibm.com | 10. Twyla Courtot
Center for Systems Management
540 N St., SW #603
Washington DC 20024-4508
phone: (202) 486-4515
email: tcourtot@csm.com |

11. John Evans
Engineous Software, Inc.
2000 Centregreen Way - Suite 100
Cary, NC 27513
phone: (919) 677-6700
email: jpevans@engineous.com
12. Eric Everton
NASA Langley Research Center
Mail Stop 125
Hampton, VA 23681
phone: (757) 864-5778
email: e.l.everton@larc.nasa.gov
13. Ponani Gopalakrishnan
IBM Research, T.J. Watson Research Center
19 Skyline Drive
Hawthorne, NY 10532
phone: (914) 784-7056
email: psg@us.ibm.com
14. Robert Graybill
DARPA/IPTO
3701 North Fairfax Drive
Arlington, VA 22203-1714
phone: (703) 696-2220
email: rgraybill@darpa.mil
15. Bernard Grossman
National Institute of Aerospace
144 Research Drive
Hampton VA 23666
phone: (757) 766-1497
email: grossman@nianet.org
16. Dana Hammond
NASA Langley Research Center
Mail Stop 125
Hampton, VA 23681
phone: (757) 864-2253
email: d.p.hammond@larc.nasa.gov
17. Thomas Hinke
NASA Advanced Supercomputing (NAS) Division, NASA Ames Research Center
Mail Stop 258-5
Moffett Field, CA 94035-1000
phone: (650) 604-3662
email: Thomas.H.Hinke@nasa.gov
18. William Jameson
HP
7505 Boxberry Terrace
Gaithersburg, MD 20879
phone: (301) 926-2370
email: William.Jameson@hp.com
19. Kennie Jones
NASA Langley Research Center
Mail Stop 125
Hampton, VA 23681
phone: (757) 864-6720
email: k.h.jones@larc.nasa.gov
20. Deb Kaplan
Fujitsu PC Corp
1800 Diagonal Road, Suite 600
Alexandria, VA 22314
phone: (703) 684-4477
email: dkaplan@fujitsupc.com
21. Peter Kazaras
Sonex Enterprises, Inc.
9990 Lee Highway
Fairfax, VA 22030
phone: (703) 691-8122, ext 230
email: pete.kazaras@sonexent.com
22. William Kneisly
IBM
6710 Rockledge Drive
Bethesda, MD 20817
phone: (301) 803-2004
email: wkneisly@us.ibm.com

23. Ken Knueven
DataSynapse, Inc.
1855 Stratford Park Place, Unit 406
Reston, VA 20190
phone: (703) 707-8398
email: ken@datasynapse.com
24. Zachary Lemnios
DARPA/IPTO
3701 North Fairfax Drive
Arlington, VA 22203
phone: (703) 696-2234
email: tedmond@snap.org
25. John Malone
NASA Langley Research Center
Mail Stop 285
Hampton, VA 23681
phone: (757) 864-8988
email: j.b.malone@larc.nasa.gov
26. Linda Marshall
Oracle Corporation
8301 Venture Drive
Waldorf, MD 20603
phone: (301) 332-8753
email: lin.marshall@oracle.com
27. Dimitri Mavris
Georgia Institute of Technology
Atlanta, GA 30332-0150
phone: (404) 894-3343
email: dimitri.mavris@ae.gatech.edu
28. Wayne McClellan
Kennedy Space Center
FL, 32899
phone: (321) 861-3704
email:
Wayne.W.McClellan@nasa.gov
29. Ed McGarr
Star Bridge Systems
7651 South Main Street
Midvale, UT 84047
phone: (801) 984-4444
email:
emcgarr@starbridgesystems.com
30. George Molloy
Marshall Space Flight Center
Huntsville, AL 35812
phone: (256) 544-2069
email: george.p.molloy@nasa.gov
31. Randy Moulic
IBM Research
1101 Kitchawan Road
Yorktown Heights, NY 10598
phone: (914) 945-1901
email: rmoulic@us.ibm.com
32. David Nelson
National Coordination Office for
Information Technology Research
and Development
4201 Wilson Blvd. Suite II-405
Arlington, VA 22230
phone: (703) 292-4873
email: nelson@itrd.gov
33. Ahmed Noor
Center for Advanced Engineering
Environments, Old Dominion
University
NASA Langley Research Center, MS
920
Hampton, VA 23681
phone: (757) 766-5233
email: a.k.noor@larc.nasa.gov
34. Michael Osias
IBM
20 Manor House Road
Budd Lae, 07828
phone: (201) 230-4905
email: mosias@us.ibm.com

35. Nickolas Paladino
Kennedy Space Center
Mail Stop YA-D6
FL 32899
phone: (321) 861-8987
email: nickolas.d.paladino@nasa.gov
36. Juliet Pao
NASA Langley Research Center
Mail Stop 124
Hampton, VA 23681
phone: (757) 864-7328
email: j.z.pao@larc.nasa.gov
37. Manish Parashar
Rutgers University
ECE, 94 Brett Road
Piscataway, NJ 08854
phone: (732) 445-5388
email: parashar@caip.rutgers.edu
38. F.G. Patterson
NASA Headquarters
300 E Street, S.W.
Washington, DC 20546
phone: (202) 358-2171
email: pat.patterson@hq.nasa.gov
39. Don Phillips
Tri Star Engineering
NASA Langley Research Center, MS 157D
Hampton, VA 23681
phone: (757) 864-4780
email: d.h.phillips@larc.nasa.gov
40. Marlon Pierce
Indiana University
501 North Morton Street
Bloomington, IL 47404
phone: (812) 856-1212
email: marpierce@indiana.edu
41. Rick Recuparo
Engineous Software, Inc.
106 Sedgwick Drive
Syracuse, NY 13203-1337
phone: (315) 428-0582
email: recuparo@engineous.com
42. Therese Rhodes
Engineous Software, Inc.
2000 Centre Green Way, Suite 100
Cary, NC 27513
phone: (919) 677-6725
email: therese.rhodes@engineous.com
43. Andrea Salas
NASA Langley Research Center
Mail Stop 159
Hampton, VA 23681
phone: (757) 864-5790
email: a.o.salas@larc.nasa.gov
44. Steve Salkeld
Platform Computing
35 Elizabeth Street, South
Brampton, Canada L6Y 1R2
phone: (905) 948-4250
email: ssalkeld@platform.com
45. Joe Schmid
Johnson Space Center
Mail Stop SD4
Houston, TX 77058
phone: (281) 483-7999
email: josef.f.schmid@nasa.gov
46. Bart Singer
NASA Langley Research Center
Mail Stop 128
Hampton, VA 23681
phone: (757) 864-2154
email: b.a.singer@larc.nasa.gov

47. Asim Smailagic
Carnegie Mellon University
5000 Forbes Avenue
Pittsburgh, PA 15213
phone: (412) 268-7863
email: asim@cs.cmu.edu
48. Maria So
Goddard Space Flight Center
Code 570
Greenbelt, MD 20771
phone: (301) 286-6113
email: maria.m.so@nasa.gov
49. Kathy Stacy
NASA Langley Research Center
Mail Stop 472
Hampton, VA 23681
phone: (757) 864-6719
email: k.stacy@larc.nasa.gov
50. Vincent Stanford
National Institute of Standards and Technology
100 Bureau Drive
Gaithersburg, MD 20878
phone: (301) 975-5399
email: stanford@nist.gov
51. Jim Steincamp
Marshall Space Flight Center
Huntsville, AL 35812
phone: (256) 544-0544
email: James W.
Steincamp@nasa.gov
52. Olaf Storaasli
NASA Langley Research Center
Mail Stop 240
Hampton, VA 23681
phone: (757) 864-2927
email: o.o.storaasli@larc.nasa.gov
53. David Tabor
Platform Computing, Inc.
phone:
email:
54. Jonas Talandis
National Center for Supercomputing Applications
University of Illinois
Urbana/Champaign
851 S. Morgan Street, Room 1120
SEO
Chicago, IL 60625
phone: (312) 996-3002
email: jonast@ncsa.uiuc.edu
55. Ric Telford
IBM Corporation
3039 Cornwallis Road
Research Triangle Park, NC 27709
phone: (919) 543-1515
email: rtelford@us.ibm.com
56. Frank Thames
NASA Langley Research Center
Mail Stop 124
Hampton, VA 23681
phone: (757) 864-5596
email: f.c.thames@larc.nasa.gov
57. Judith Utley
AMTI/NASA Ames Research Center
phone:
email: utley@nas.nasa.gov
58. Tamer Wasfy
Advanced Science and Automation Corporation
20170 E. Magnolia Court
Smithfield, VA 23430
phone: (757) 469-6839
email: tamer@ascience.com

59. Jim Yardley
Star Bridge Systems
7651 South Main
Midvale, Utah 84047
phone: (801) 984-4444
email:
jyardley@starbridgesystems.com
60. Youngjin Yoo
Case Western Reserve University
10900 Euclid Avenue
Cleveland, OH 44106
phone: (216) 368-0790
email: yxy23@po.cwru.edu
61. Victor Zue
Massachusetts Institute of
Technology
200 Technology Square
Cambridge, MA 02139
phone: (617) 258-6206
email: zue@lcs.mit.edu

**Perspectives on
Emerging/Novel Computing Paradigms and Future Aerospace
Workforce Environments**

Ahmed K. Noor
Center for Advanced Engineering Environments
Old Dominion University
NASA Langley Research Center
Hampton, VA

INTRODUCTION

The accelerating pace of the computing technology development shows no signs of abating. Computing power reaching 100 Tflop/s is likely to be reached by 2004 and Pflop/s (10^{15} Flop/s) by 2007. The fundamental physical limits of computation, including information storage limits, communication limits and computation rate limits will likely be reached by the middle of the present millennium. To overcome these limits, novel technologies and new computing paradigms will be developed.

An attempt is made in this overview to put the diverse activities related to new computing-paradigms in perspective and to set the stage for the succeeding presentations. The presentation is divided into five parts (Figure 1). In the first part, a brief historical account is given of development of computer and networking technologies. The second part provides brief overviews of the three emerging computing paradigms – grid, ubiquitous and autonomic computing. The third part lists future computing alternatives and the characteristics of future computing environment. The fourth part describes future aerospace workforce research, learning and design environments. The fifth part lists the objectives of the workshop and some of the sources of information on future computing paradigms.

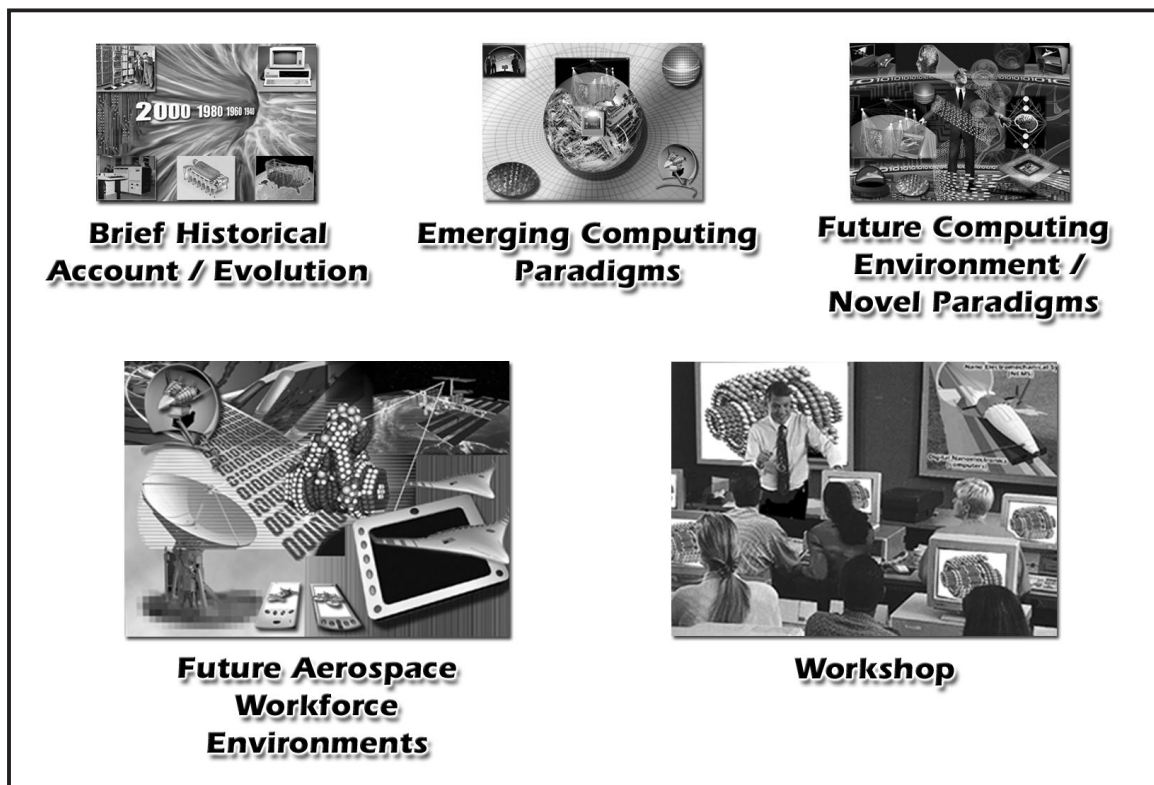


Figure 1

BRIEF HISTORICAL ACCOUNT OF THE DEVELOPMENT OF COMPUTER AND NETWORKING TECHNOLOGIES

The field of computing is less than sixty years old. The first electronic computers were built in the 1940s as part of the war effort. The first transistor was invented in 1947. By 1950s, IBM and Univac built business computers, intended for scientific and mathematical calculations to determine ballistic trajectories and break ciphers. Soon other companies joined the effort – names like RCA, Burroughs, ICL and General Electric – most of whom disappeared or left the computer business. The first programming languages – Algol, FORTRAN, Cobol, and Lisp – were designed in the late 1950s, and the first operating system in the early 1960s. The first computer chip appeared in the late 1970s, the personal computer around the same time, and the IBM PC in 1981. Ethernet was invented in 1973 and did not appear in the market until 1980. It operated at 10 megabits per second (10 Mb/s) and increased to 1 Gb/s (10^9 bits/s) in 1997. The internet descended from the ARPANET in 1970s, and the World Wide Web was created in 1989 (see Figure 2).

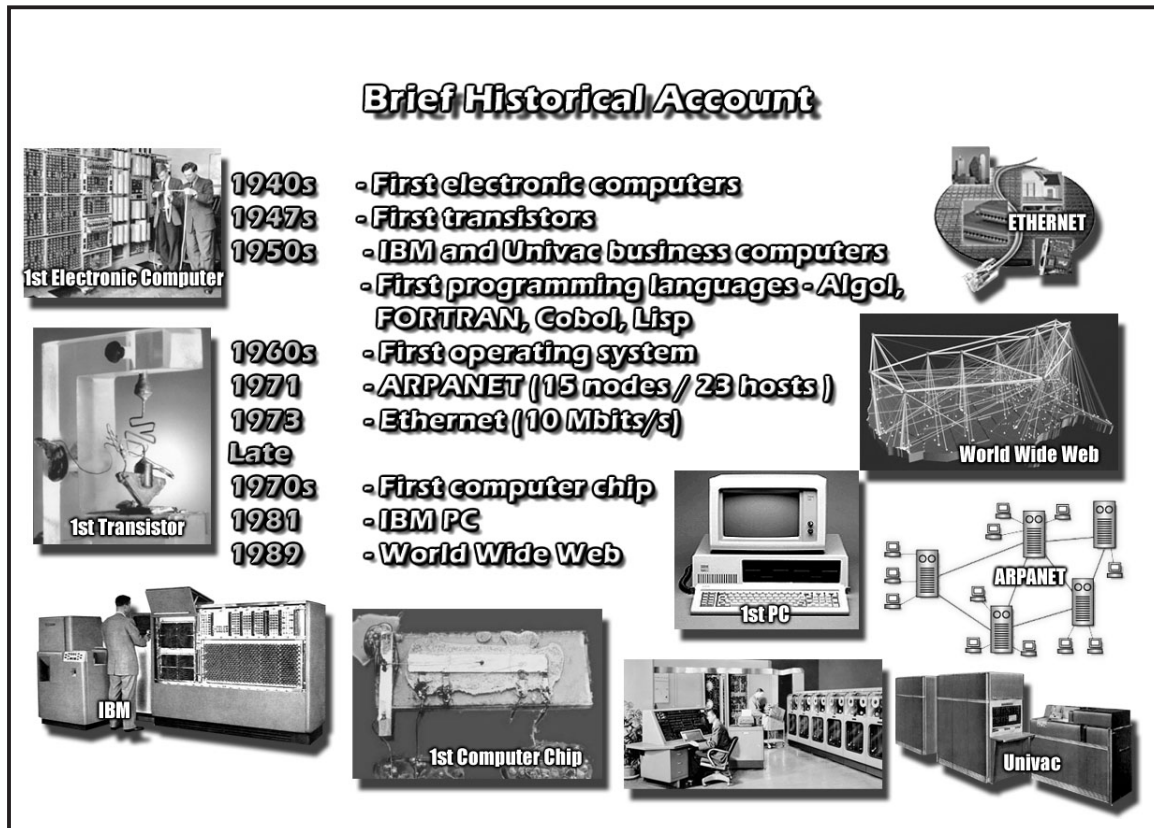


Figure 2

EVOLUTION OF MICROPROCESSORS

Although the first computers used relays and vacuum tubes for the switching elements, the age of digital electronics is usually said to have begun in 1947, when a research team at Bell Laboratories designed the first transistor. The transistor soon displaced the vacuum tube as the basic switching element in digital design. The nerve center for a computer, or a computing device, is its integrated circuit (IC or chip), the small electronic device made out of a semiconductor material. Integrated circuits, which appeared in the mid-1960's and allowed mass fabrication of transistors on silicon substrates are often classified by the number of transistors and other electronic components they contain. The ever-increasing number of devices packaged on a chip has given rise to the acronyms SSI, MSI, LSI, VLSI, ULSI, and GSI, which stand for small scale (1960s – with up to 20 gates per chip), medium-scale (late 1960's – 20-200 gates), large-scale (1970s – 200-5000 gates per chip), very large-scale (1980s – over 5000 gates per chip), ultra large-scale (1990s – over million transistors per chip), and giga-scale integration (over billion transistors per chip), respectively (Figure 3).

In 1965, Gordon Moore hypothesized that processing power (number of transistors and computing speed) of computer chips was doubling every 18-24 months or so. For nearly four decades the chip industry has marched in lock step to this pattern or rule of thumb, which is referred to as Moore's law (see Figure 3).

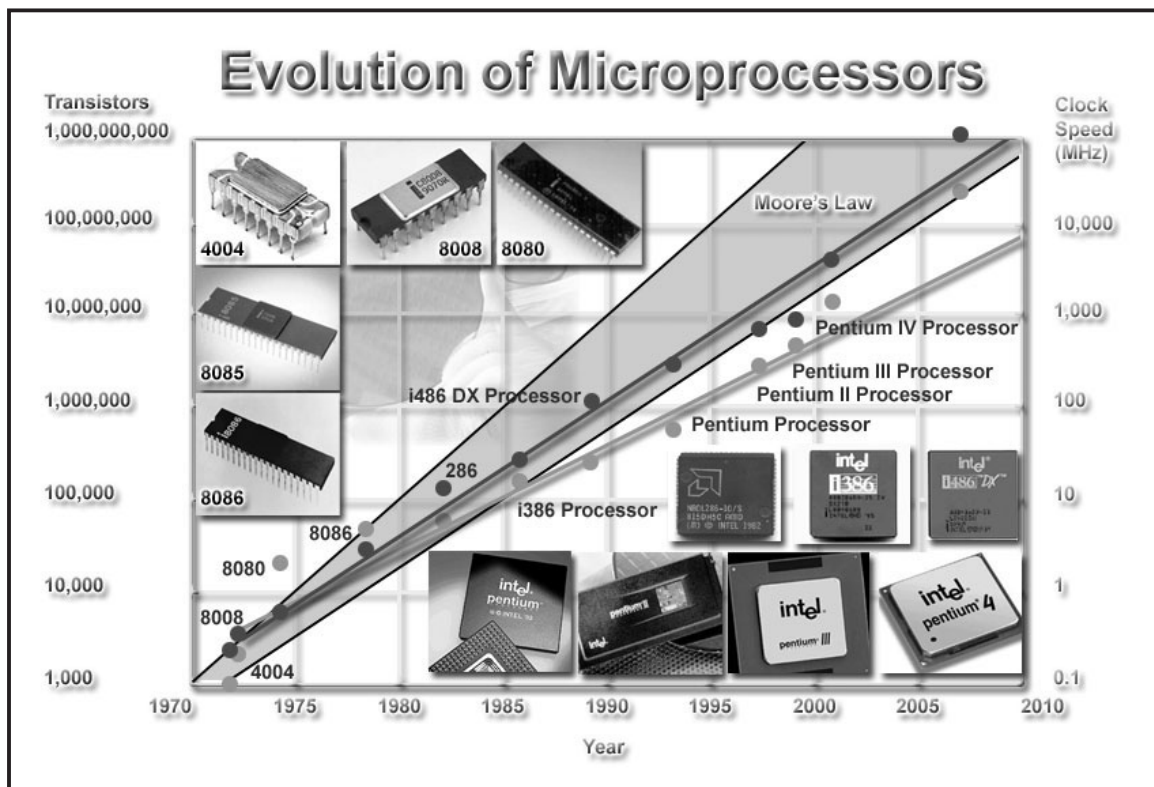


Figure 3

GROWTH IN COMPUTER SPEED AND SHIFT IN HARDWARE TECHNOLOGY

Advances in microprocessor technology resulted in increasing the speed of computers by more than trillion times during the last five decades, while dramatically reducing the cost (Figure 4).

A number of technologies have been used to achieve ultra fast logic circuits. These include use of: new material systems such as gallium arsenide (Ga As); multichip modules (MCM); monolithic and hybrid wafer-scale integration (WSI); new transistor structures such as the quantum-coupled devices using hetero-junction-based super lattices; and optical interconnections and integrated optical circuits. More recently, the use of carbon nanotubes as transistors in chips; clockless (asynchronous) chips and; hyper-threading, which makes a single CPU act in some ways like two chips, have been demonstrated.

The incessant demand for computing power to enable accurate simulation of complex phenomena in science and engineering has resulted in the development of a class of general-purpose supersystems designed for extremely high-performance throughput, and new paradigms for achieving the high-performance. These include:

- Vector/pipeline processing
- Parallel processing on multiple (hundreds or thousands) CPUs, and
- Multitasking with cache memory microprocessors

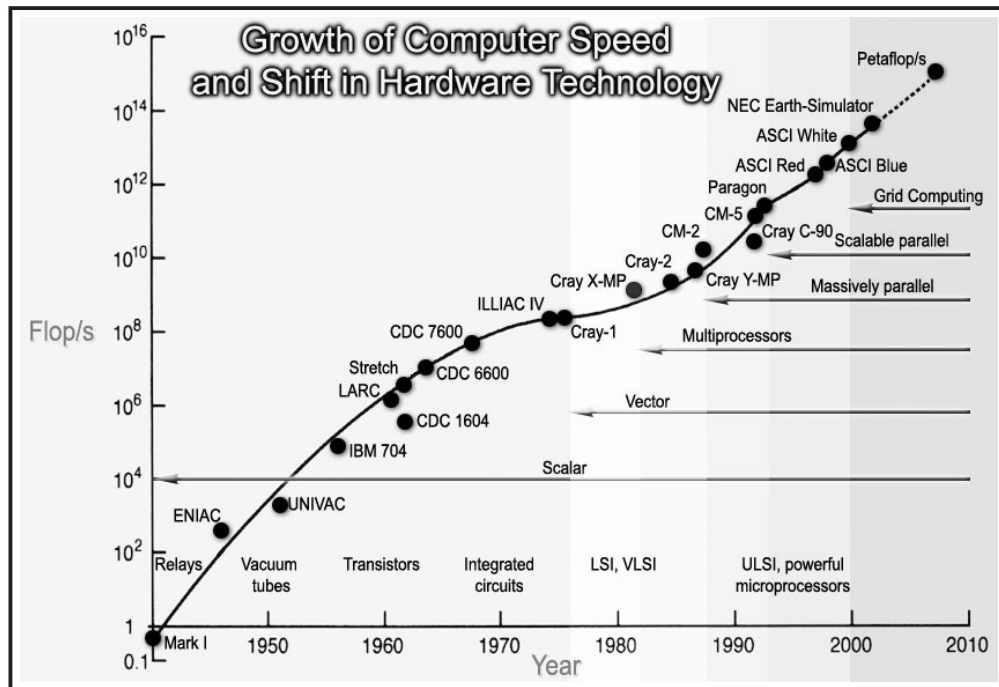


Figure 4

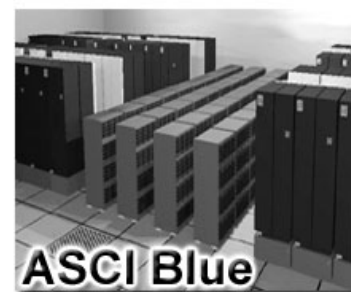


Figure 4 (continued)

TOP FIVE SUPERCOMPUTER SITES

Although the peak performance of the first generation supersystems was less than 100 Mflop/s, the gigaflop barrier (1 Gflop/s) was passed in 1988/89, and the teraflop barrier (1 Tflop/s) in 1996/7. In 1995, the US Department of Energy supported the development of three terascale machines through its Accelerated Strategic Computing Initiative (ASCI). The three machines are: ASCI Red, with 9,472 Intel Pentium II Xeon processors – 2.379 Tflop/s at Sandia National Labs; ASCI Blue Mountain with 5,856 IBM PowerPC 604E processors – 1.608 Tflop/s at Los Alamos National Labs; and ASCI White with 8,192 IBM Power 3-II processors – 7.226 Tflop/s at Lawrence Livermore National Lab.

To date, there are over 17 terascale machines worldwide. The maximum performance reported today is 35.86 Tflop/s of the Earth Simulator at Kanazawa, Japan, which consists of 5,104 vector processors (with peak performance of 40 Tflop/s).

The top five supercomputer sites, based on the Linpack benchmark are shown in Figure 5.

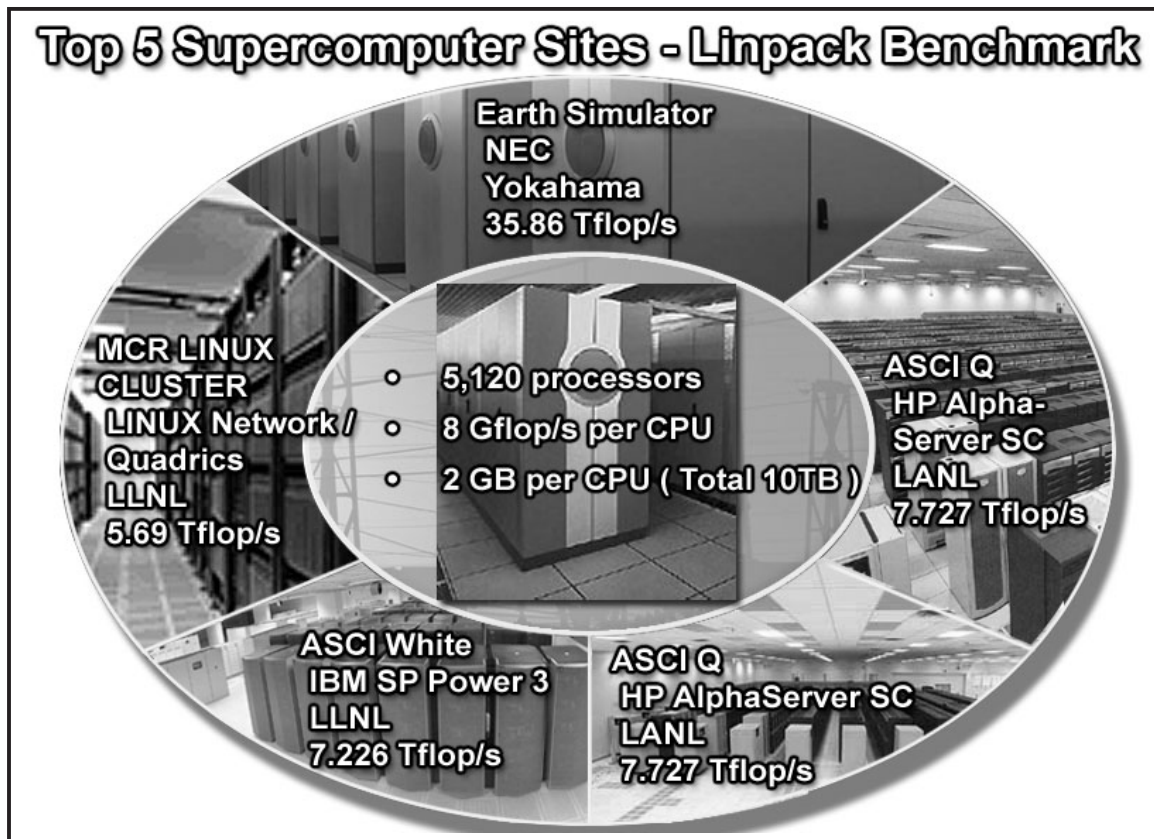


Figure 5

PETAFLOP SUPERSYSTEM

In December 1999, IBM announced a five year effort to build a petaflop (10^{15} Flop/s) supersystem – The Blue Gene Project. The project has the two primary goals of advancing the state of the art of biomolecular simulation, and computer design of extremely large-scale systems. Two systems are planned: Blue Gene/L, in collaboration with Lawrence Livermore National Lab, which leverages high speed interconnect and system-on-a-chip technologies and has a peak performance of 200 Tflop/s; and Blue Gene/P, the petaflop-scale system. The system will consist of more than one million processors, each capable of one billion operations per second. Thirty-two of these ultra-fast processors will be placed on single chip (32 Gflop/s). A compact two-foot by two-foot board containing 64 of these chips will be capable of 2 Tflop/s. Eight of the boards will be placed in 6-foot high racks (16 Tflop/s) and the final system will consist of 64 racks linked together to achieve the one Pflop/s performance (Figure 6).

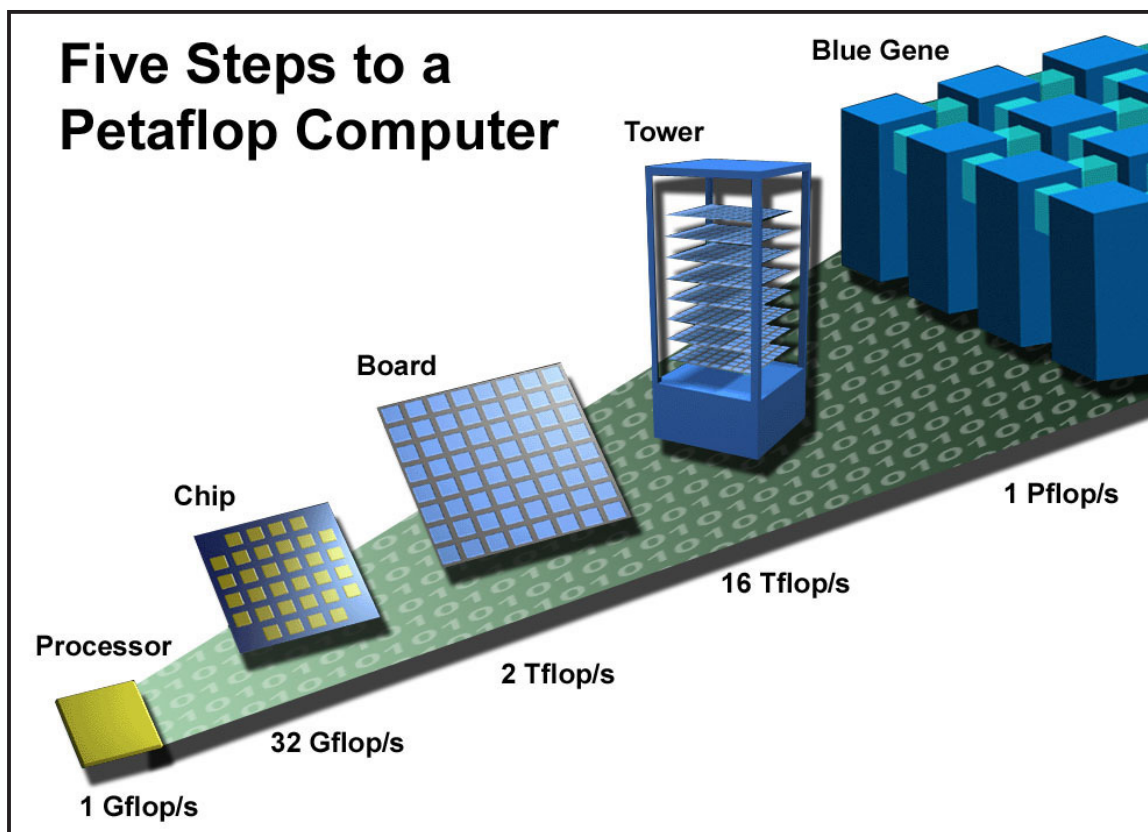


Figure 6

EVOLUTION OF HUMAN-COMPUTER INTERFACES

Figure 7 shows the evolution of human-computer interfaces. During the period of 1940's through 1970's, static interfaces for main frames were used in the form of teletype style. This was followed in the 1980's by more flexible interfaces for PCs – Windows, mouse and graphical tablet. With many computing devices available for single users, adaptive interfaces with more functionality and communication became available. The emergence of grid/pervasive computing paradigms is providing an impetus for intelligent neural, perceptual, attentive and other advanced interfaces which integrate adaptive interfaces with intelligent agents for making intelligent help and tutoring available to the user.

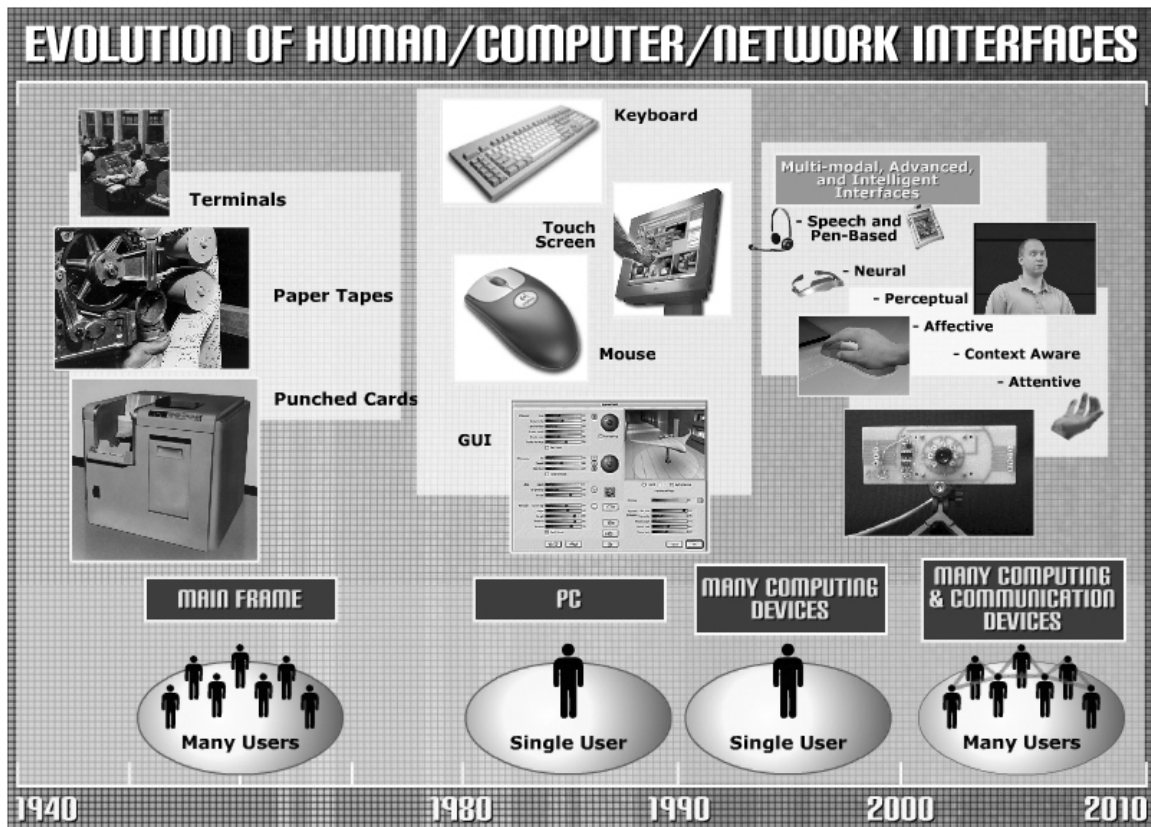


Figure 7

EMERGING COMPUTING PARADIGMS

The rapidly increasing power of computers and networks, and the trend of computers getting smaller, along with the increasing complexity of computing systems and the associated cost to manage them, led to three emerging computing paradigms, namely (Figure 8),

- Grid Computing,
- Ubiquitous/Pervasive Computing, and,
- Autonomic Computing

The three paradigms are described subsequently.

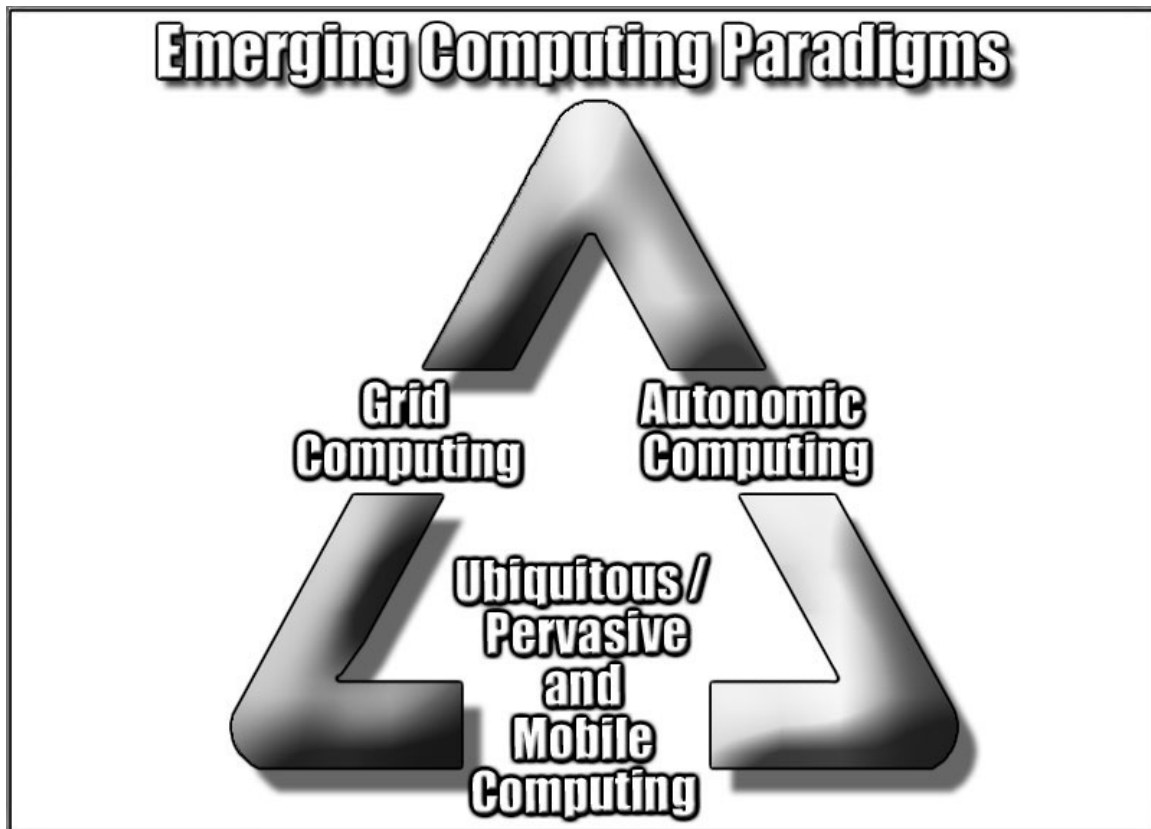


Figure 8

GRID COMPUTING

The rapidly increasing power of computers and networks in the 1990s, led the new paradigm of distributed computing. A flurry of experiments were conducted on “peer-to-peer” computing, all devoted to harnessing the computer power and storage capacity of idle desktop machines. These included cluster computing - using networks of standard single-processor workstations to solve single problems. At the same time, the high-performance computer community began the more ambitious experiments in *metacomputing*. The objective of Metacomputing was to make many distributed computers function like one giant computer – metasystem (e.g., the virtual national machine). Metasystems give users the illusion that the files, databases, computers and external devices they can reach over a network constitute one giant transparent computational environment.

The term *grid computing* is now used to refer to massive integration of computer systems to offer performance unattainable by a single machine. It provides pervasive, dependable, consistent, and inexpensive access to facilities and services that live in cyberspace, assembling and reassembling them on the fly to meet specified needs (Figure 9).

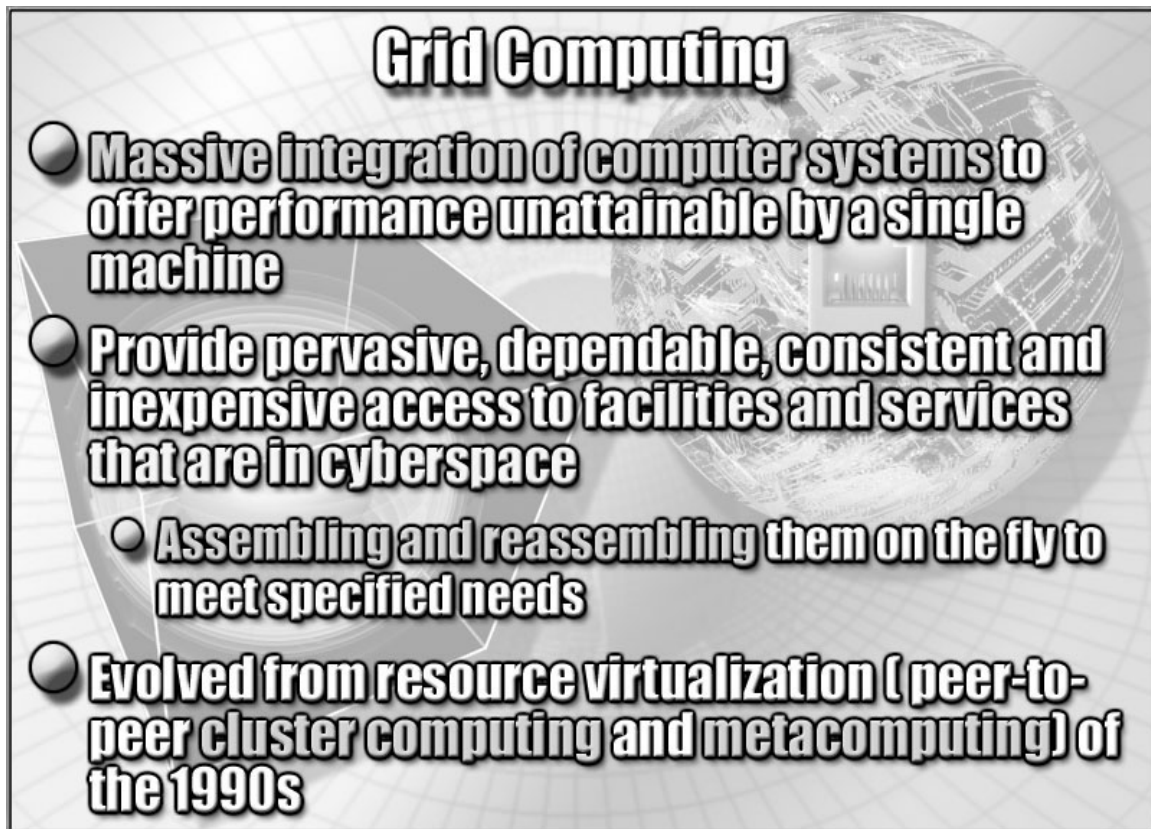


Figure 9

GRID TECHNOLOGIES AND INFRASTRUCTURE

The essential building blocks of grid computing are: Fast processors, parallel computer architectures, advanced optical networks, communication protocols, distributed software structures and security mechanisms (Figure 10).

Grid technologies enable the clustering of a wide variety of geographically distributed resources, such as high-performance computers, storage systems, data sources, special devices and services that can be used as a unified resource.

Although grid technologies are currently distinct from other major technology trends, such as internet, enterprise, distributed, and peer-to-peer computing, these other trends can benefit significantly from growing into the problem spaces addressed by grid technologies.

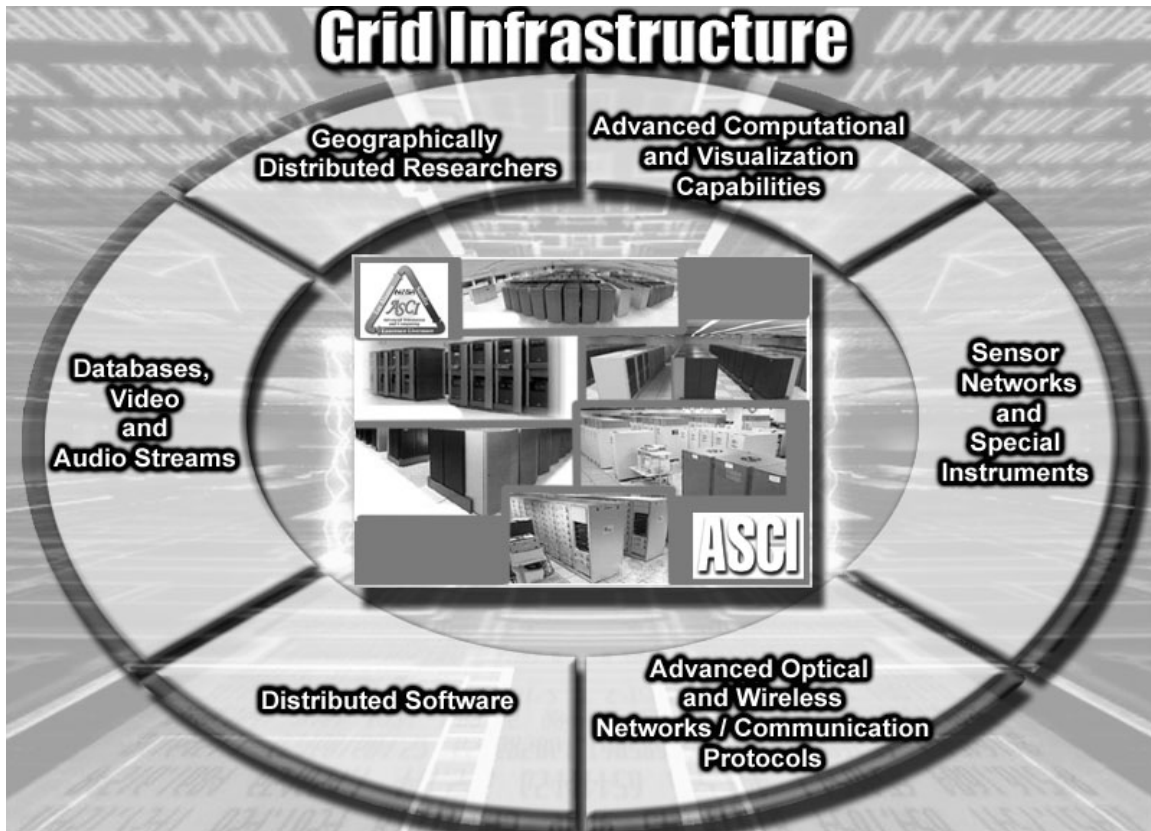


Figure 10

GRID COMPUTING PROJECTS

Once the concept of grid computing was introduced, several grid computing projects were launched all over the world. A sampling of grid computing projects are listed in Figure 11. In the future, grids of every size will be interlinked. The “supernodes” like TeraGrid will be networked clusters of supersystems serving users on a national or international scale. Still more numerous will be the millions of individual nodes: personal machines that users plug into the grid to tap its power as needed. With wireless networks and miniaturization of components, that can evolve into billions of sensors, actuators and embedded processors as micronodes.

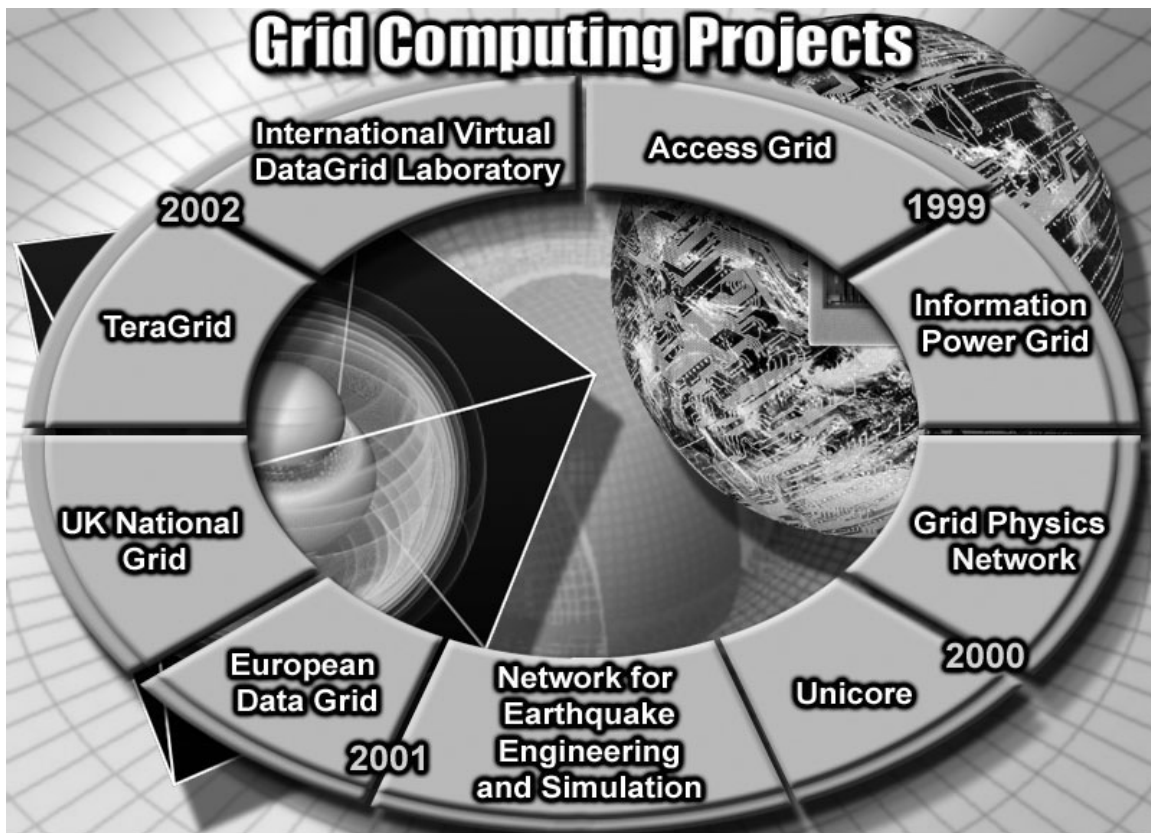


Figure 11

UBIQUITOUS / PERVASIVE COMPUTING

The trend of computers getting smaller is likely to lead to an environment with computing functionality embedded in physical devices that are widely distributed and connected in a wireless web.

In a seminal article written in 1991, Mark Weiser described a hypothetical world in which humans and computers were seamlessly united. This vision was referred to as *ubiquitous computing*. Its essence was the creation of environment saturated with computing and communication, yet gracefully integrated with human users.

In the mid-1990s, the term pervasive computing came to represent the same vision as that described by Weiser.

The key components of ubiquitous/pervasive computing are (Figure 12):

- *Pervasive devices*, including:
 - Small, low-powered hardware (CPU, storage, display devices, sensors),
 - Devices that come in different sizes for different purposes, and
 - Devices that are aware of their environment, their users, and their locations,
- *Pervasive communication* – a high degree of communication among devices, sensors and users provided by ubiquitous and secure network infrastructure (wireless and wired) and mobile computing,
- *Pervasive interaction* – more natural and human modes of interacting with information technology, and
- *Flexible, adaptable distributed systems* – dynamic configuration, functionality on demand, mobile agents and mobile resources

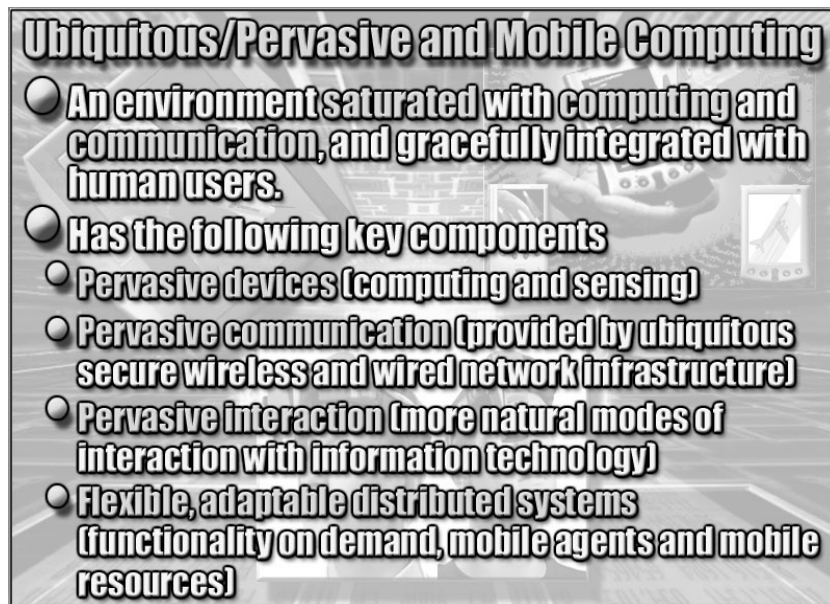


Figure 12

PERVASIVE COMPUTING FRAMEWORK

The technological advances necessary to build a pervasive computing environment fall into four broad areas (Figure 13): devices, networking, middleware and applications. Middleware mediates interactions with the networking kernel on the user's behalf and keeps users immersed in the pervasive computing space. The middleware consists mostly of firmware and software

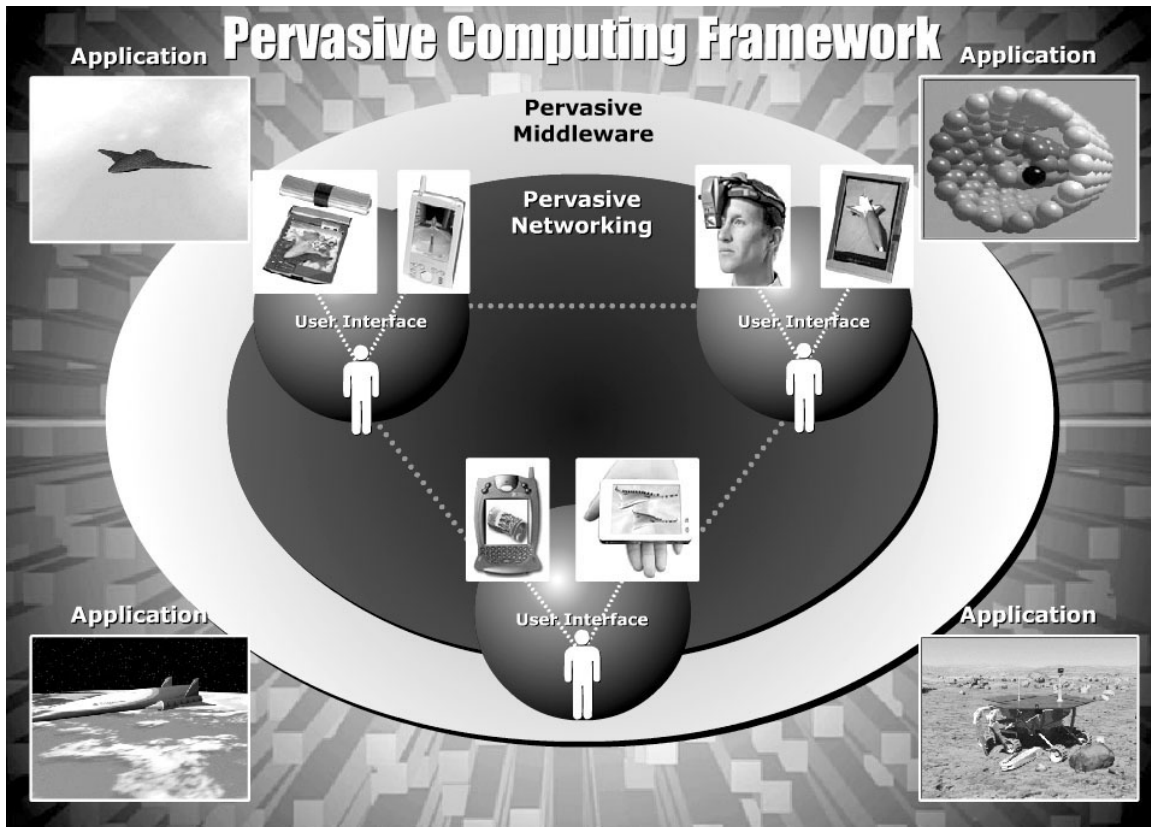


Figure 13

PERVASIVE COMPUTING INITIATIVES

A list of some of the pervasive computing initiatives is given in Figure 14. These include university initiative (AURA of Carnegie Mellon University, Endeavor of the University of California at Berkeley, the Oxygen Project of MIT, and Portolano Project of the University of Washington); Industry/university initiatives (Sentient Computing, a joint project of AT&T Laboratories and Cambridge University in the UK); and industry projects (Cooltown of Hewlett-Packard, EasyLiving of Microsoft Research Vision Group and WebSphere Everyplace of IBM).

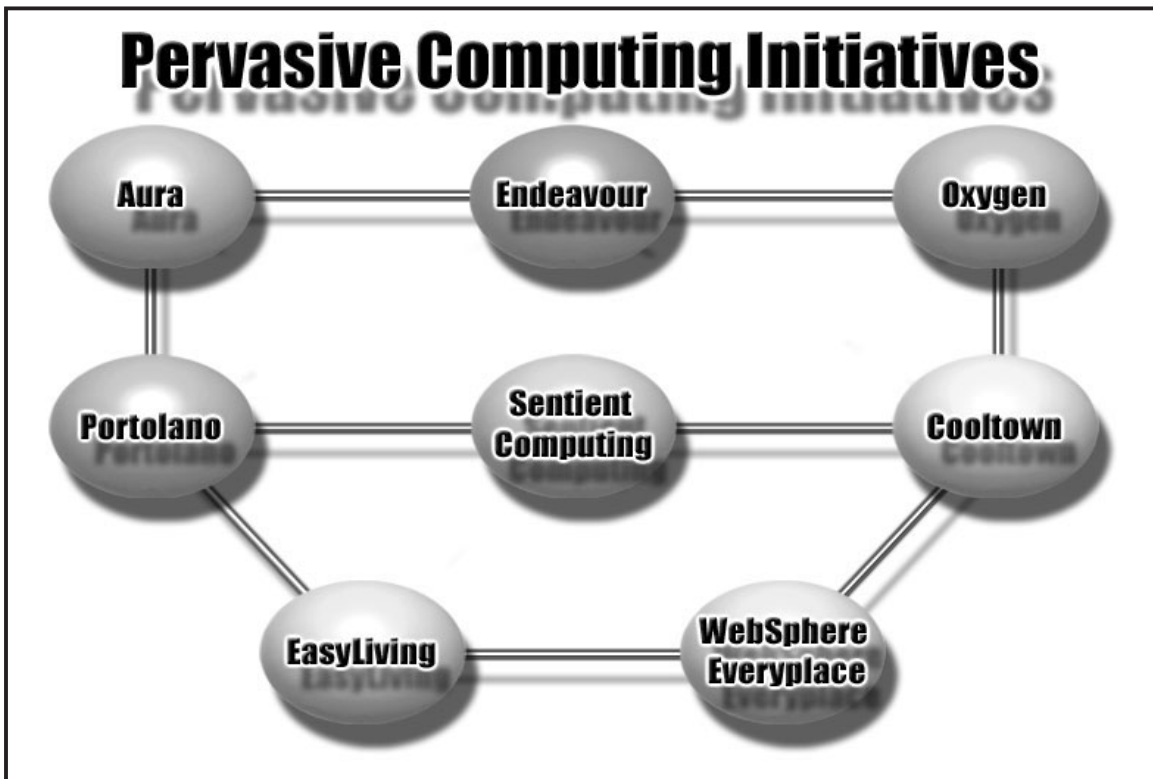


Figure 14

AUTONOMIC COMPUTING

The increasing capacity and complexity of the emerging computing systems, and the associated cost to manage them, combined with a shortage of skilled workforce are providing the motivation for a paradigm shift to systems that are self-managing, self-optimizing, and do not require the expensive management services needed today. A useful biological metaphor is found in the autonomic nervous system of the human body – it tells the heart how many times to beat, monitors the body temperature, and adjusts the blood flow, but most significantly, it does all this without any conscious recognition or effort on the part of the person - hence the name *autonomic computing* was coined.

Autonomic computing is a new research area led by IBM focusing on making computing systems smarter and easier to administer. Many of its concepts are modelled on self-regulating biological systems.

Autonomic computing is envisioned to include the ability of the system to respond to problems, repair faults and recover from system outages without the need for human intervention. An autonomic computing system consists of a large collection of computing engines, storage devices, visualization facilities, operating systems, middleware and application software (Figure 15).

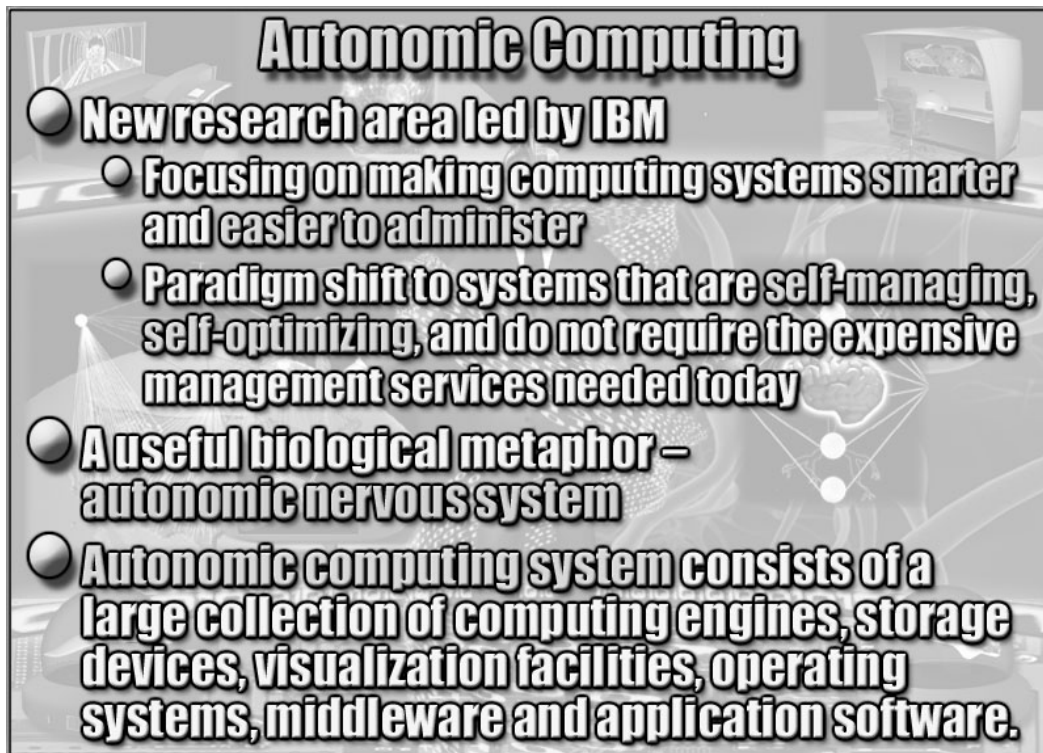


Figure 15

CHARACTERISTICS OF AUTONOMIC COMPUTING

Autonomic computing is envisioned to combine the following seven characteristics (Figure 16):

1. *Self-defining* - Has detailed knowledge of its components, current status, ultimate capacity and performance, and all connections to other systems.
2. *Self-configuring* – can configure and reconfigure itself under varying and unpredictable conditions. System configuration or setup must occur automatically, as must dynamic adjustments to that configuration to handle changing environments.
3. *Self-optimizing* – never settles for status quo. Always looks for ways to optimize its performance. Monitors constituent parts, and metrics, using advanced feedback control mechanisms and makes changes (e.g., fine-tune workflow) to achieve predetermined system goals.
4. *Self-healing* – able to recover from routine and extraordinary events that might cause some components to malfunction or damage. It must be able to discover problems, reconfigure the system to keep functioning smoothly.
5. *Self-protecting* – detect, identify and protect itself against various types of failure. Maintains overall system security and integrity.
6. *Contextually Aware* – This is almost self-optimization turned outward. The system must know the environment and the context of the surrounding activity, and adapts itself (in real-time) accordingly.
7. *Anticipatory* – anticipates the optimized resources, configuration, and components needed.

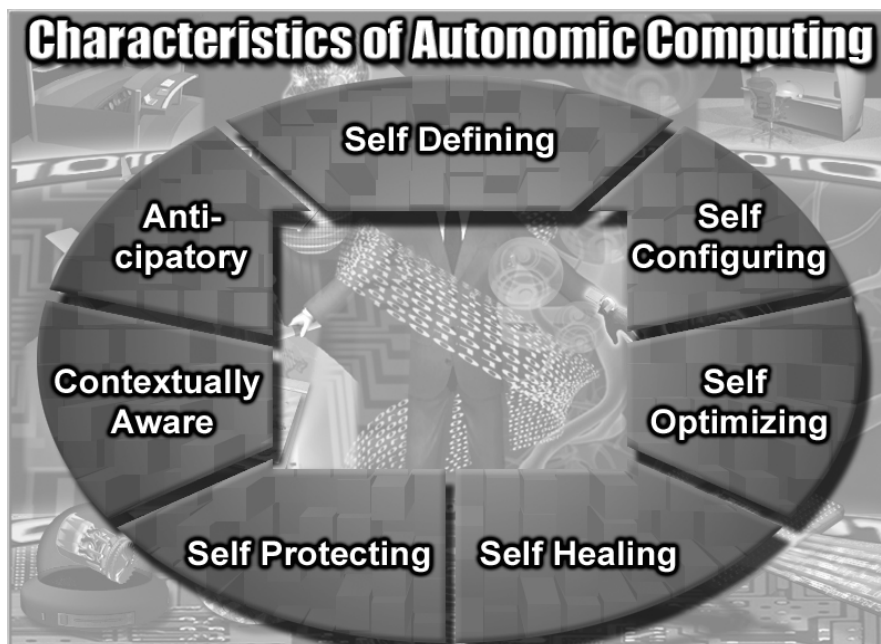


Figure 16

FUTURE COMPUTING ALTERNATIVES

Silicon-based technology is expected to reach its physical limits in the next decades. But silicon and computing are not inextricably linked, although they often seem to be. For example, when silicon microelectronics reaches ultimate physical limits a number of new approaches and technologies have already been proposed. These include (Figure 17):

- Quantum computing,
- Molecular computing,
- Chemical and biochemical computing,
- DNA computing, and
- Optical and optoelectronic computing

None of these approaches is ready to serve as an all-purpose replacement for silicon. In fact, some approaches may be only appropriate as specialized methods in particular niches, such as high-level cryptography.

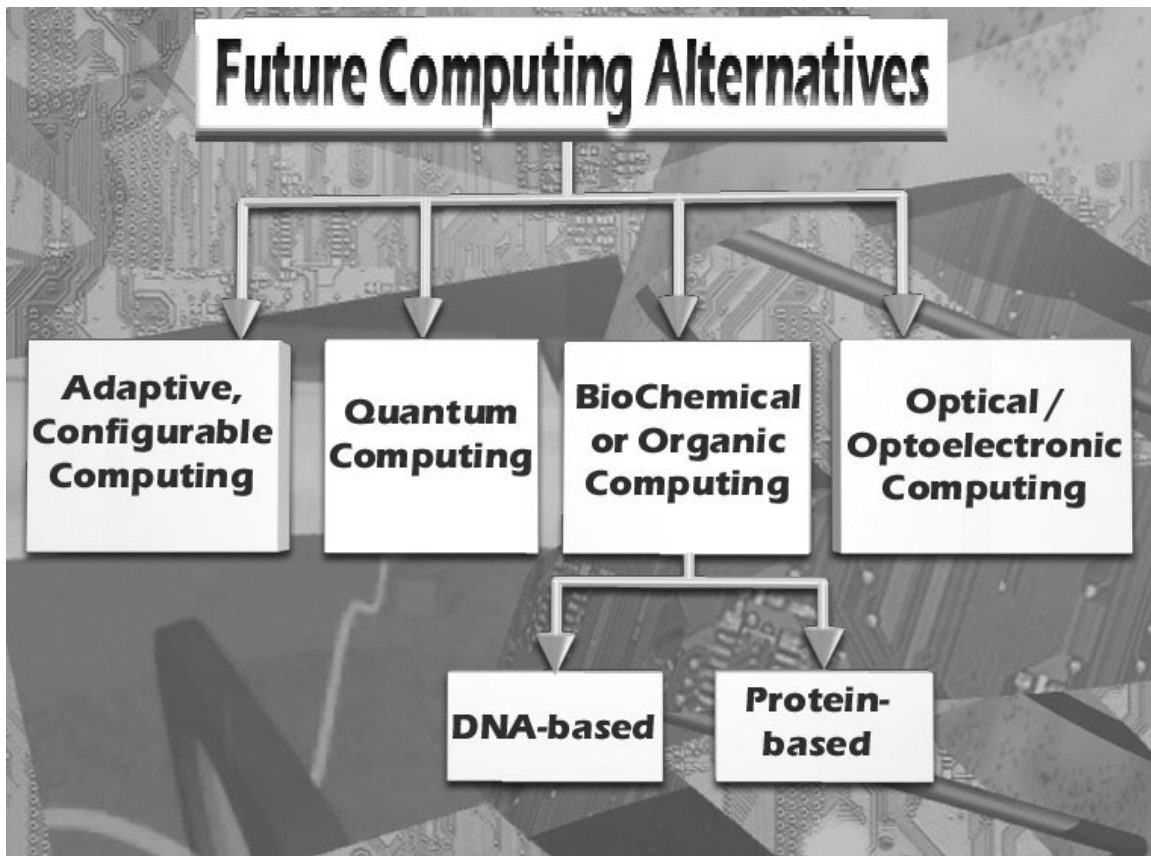


Figure 17

FUTURE COMPUTING ENVIRONMENT

Significant advances continue to be made in the entire spectrum of computing and communication technologies. Speculations about the future of computers and computing have been attempted in several monographs. Herein, only the emerging trends are identified, which include (Figure 18):

- An evolving computing paradigm combining ubiquitous / mobile / cognitive / autonomic computing and including:
 - Smart, self-regulating computing systems covering a spectrum of handheld, embedded and wearable information appliances and devices
 - Wide range of devices to sense, influence and control the physical world
 - Optical networks supplement by wireless communication
- Human-computer symbiosis characterized by:
 - Natural cooperative human-machine collaboration
 - Intelligent affective technologies to allow computers to know user's emotional states
 - Humans, sensors and computing devices seamlessly united
- Hierarchical knowledge nets:
 - Computer-supported distributed collaboration
 - Augmented / mixed reality and tele-immersion facilities
 - Advanced modeling, simulation and multisensory visualization

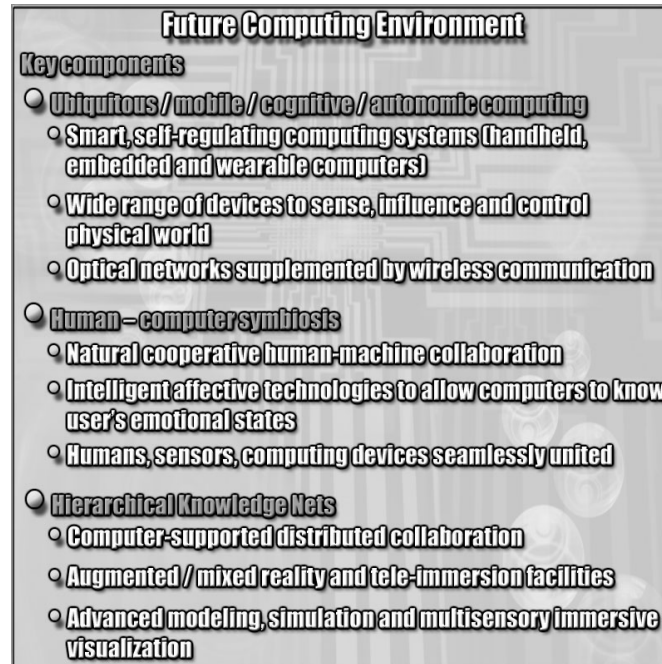


Figure 18

EXAMPLES OF FUTURE AEROSPACE SYSTEMS AND SOME OF THEIR CHARACTERISTICS

The realization of NASA's ambitious goals in aeronautics and space with the current national budget constraints will require new kinds of aerospace systems and missions that use novel technologies and manage risk in new ways. Future aerospace systems must be autonomous, evolvable, resilient, and highly distributed. Two examples are given in Figure 19. The first is a biologically inspired aircraft with self-healing wings that flex and react like living organisms. It is built of a multifunctional material with fully integrated sensing and actuation, and unprecedented levels of aerodynamic efficiencies and aircraft control. The second is an integrated human-robotic outpost, with biologically inspired robots. The robots could enhance the astronaut's capabilities to do large-scale mapping, detailed exploration of regions of interest, and automated sampling of rocks and soil. They could enhance the safety of the astronauts by alerting them to mistakes before they are made, and letting them know when they are showing signs of fatigue, even if they are not aware of it.



Figure 19

ENABLING TECHNOLOGIES FOR FUTURE AEROSPACE SYSTEMS

The characteristics of future aerospace systems identified in Figure 18 are highly coupled and require the synergistic coupling of the revolutionary and other leading-edge technologies listed in Figure 20. The four revolutionary technologies are nanotechnology, biotechnology, information/knowledge technology, and cognitive systems technology. The other leading-edge technologies are high-productivity computing; high-capacity communication; multiscale modeling, simulation and visualization; virtual product development; intelligent software agents; reliability and risk management; human performance, and human-computer symbiosis.

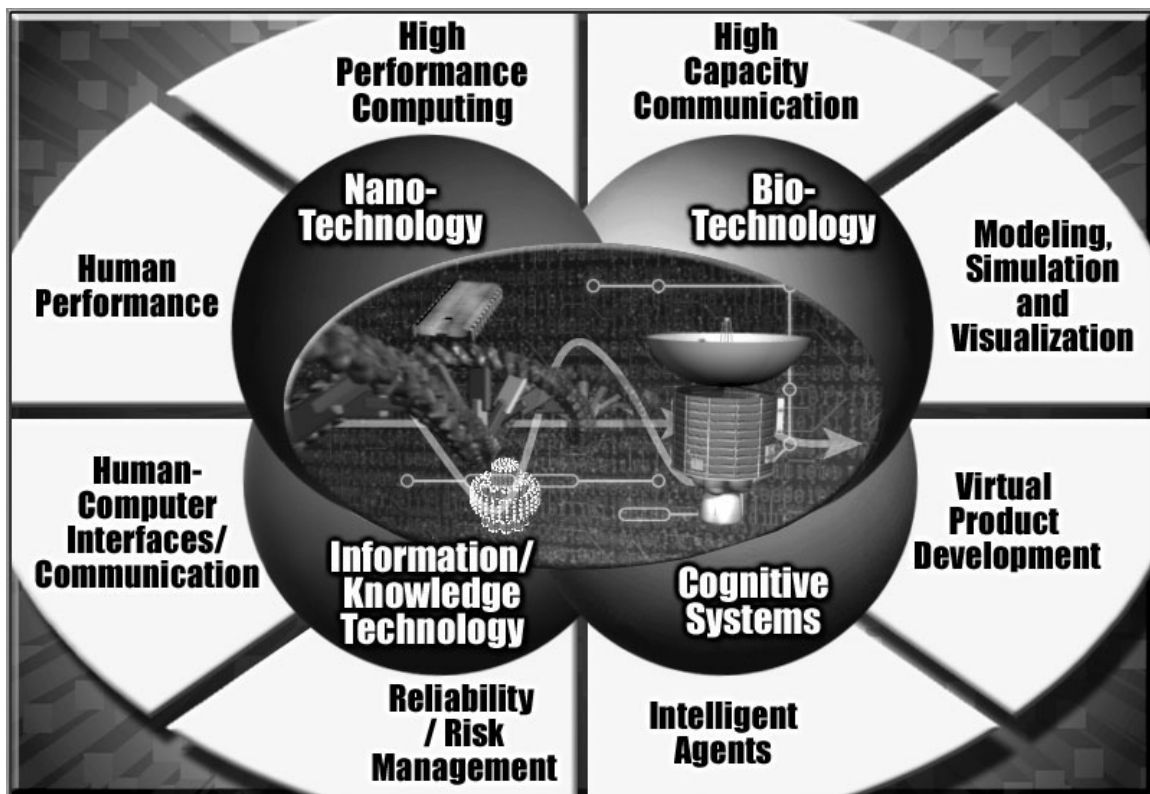


Figure 20

THREE NASA INITIATIVES

The realization of NASA's ambitious goals will require a diverse, technically skilled workforce – a new generation of scientists and engineers who can work across traditional disciplines and perform in a rapidly changing environment.

NASA has developed a number of new initiatives for assured workforce development. These include University Research, Engineering, and Technology Institutes (URETIs), the National Institute of Aerospace (NIA), and the Hierarchical Research and Learning Network (HRLN) (see Figure 21). The overall goal of these activities is to strengthen NASA's ties to the academic community through long-term sustained investment in areas of innovative and long-term technology critical to future aerospace systems and missions. At the same time, the three activities will enhance and broaden the capability of the nation's universities to meet the needs of NASA's science and technology programs.



Figure 21

HIERARCHICAL RESEARCH AND LEARNING NETWORK

The Hierarchical Research and Learning Network (HRLN) is a pathfinder project for the future aerospace workforce development. It aims at creating knowledge organizations in revolutionary technology areas which enable collective intelligence, innovation and creativity to bear on the increasing complexity of future aerospace systems. This is accomplished by building research and learning networks linking diverse interdisciplinary teams from NASA and other government agencies with universities, industry, technology providers, and professional societies (Figure 22) in each of the revolutionary technology areas and integrating them into the HRLN.

HRLN is envisioned as a neural network of networks. It is being developed by eight university teams, led by Old Dominion University's Center for Advanced Engineering Environments.



Figure 22

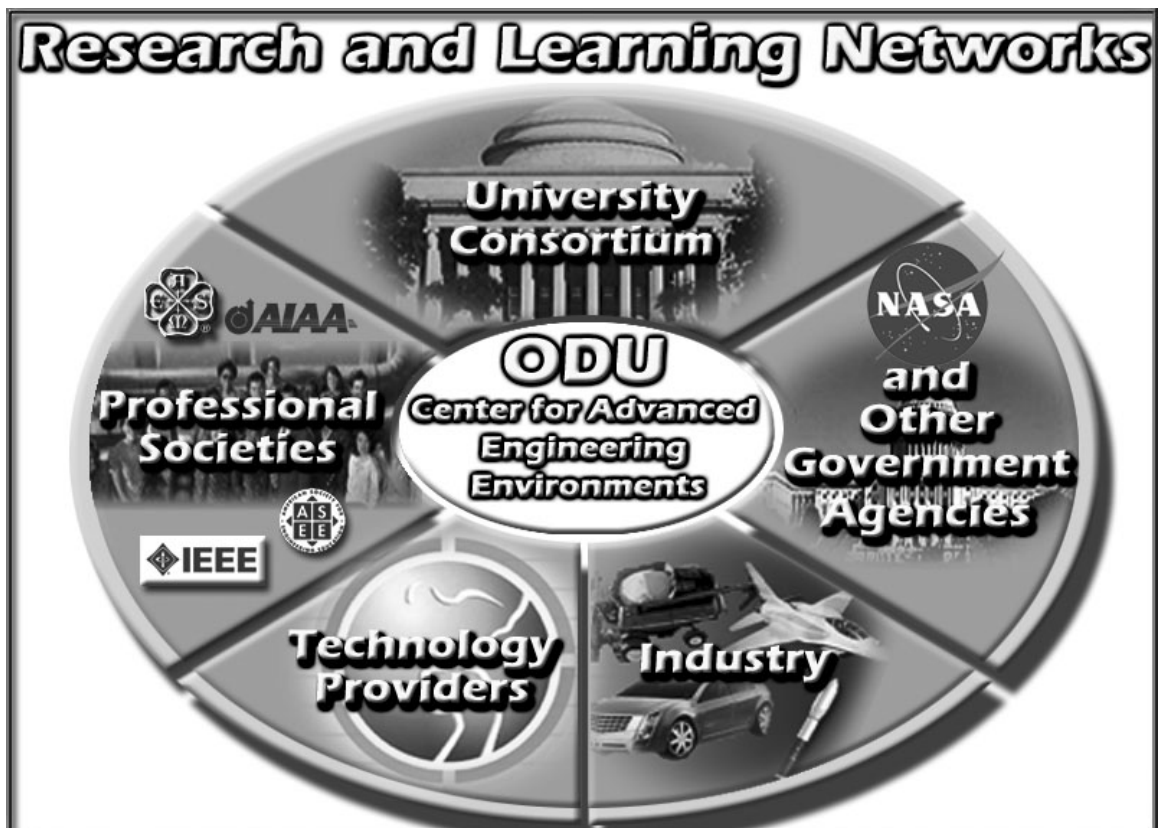


Figure 22 (continued)

IMPLEMENTAION OF HIERARCHICAL RESEARCH AND LEARNING NETWORK

The phases of implementing HRLN are shown in Figure 23. The first phase involves development of learning modules and interactive virtual classrooms in revolutionary technology areas, simulators of unique test facilities at NASA, and a telescience system – an online multi-site lab that allows real-time exchange of information and remote operation of instrumentation by geographically distributed teams. These facilities will be integrated into adaptive web learning portals in the second phase, which evolve into robust learning networks. In the final phase, the learning networks are integrated into the HRLN.

Implementation of Hierarchical Research and Learning Network

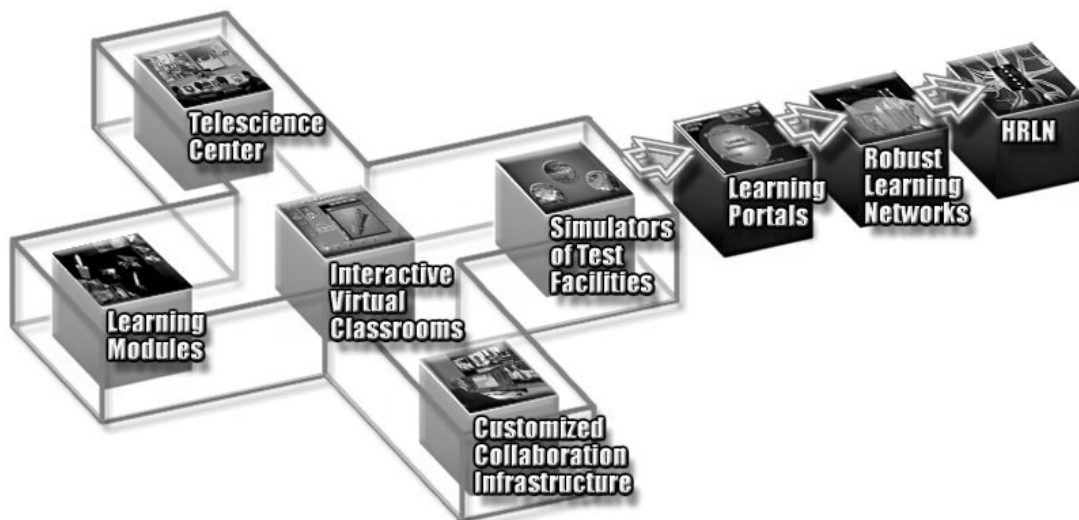


Figure 23

ADAPTIVE WEB LEARNING PORTAL

The Adaptive Web Learning Portal being developed as part of the HRLN project has the following major components (Figure 24):

- Advanced multimodal interfaces,
- Knowledge repository,
- Blended learning environment incorporating the three environments: expert-managed, self-paced, and collaborative,
- Learning management system, and
- Customized collaboration infrastructure

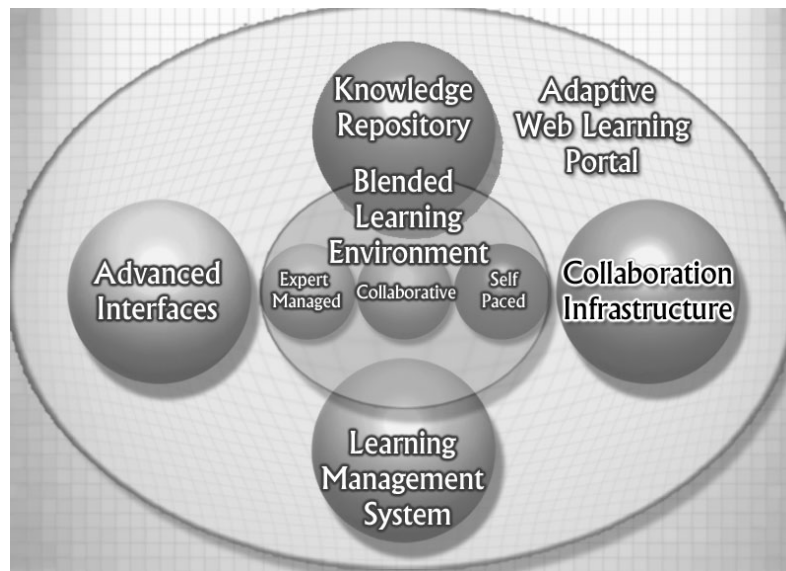


Figure 24

INTELLIGENT DESIGN ENVIRONMENT

The future design environment will enable collaborative distributed synthesis to be performed by geographically dispersed interdisciplinary/multidisciplinary teams. It will include flexible and dynamic roomware (active spaces/collaboration landscape) facilities consisting of (Figure 25):

- Portable and stationary information devices
- Novel multiuser smart displays
- Telepresence and other distributed collaboration facilities
- Novel forms of multimodal human/network interfaces
- Middleware infrastructures and intelligent software agents

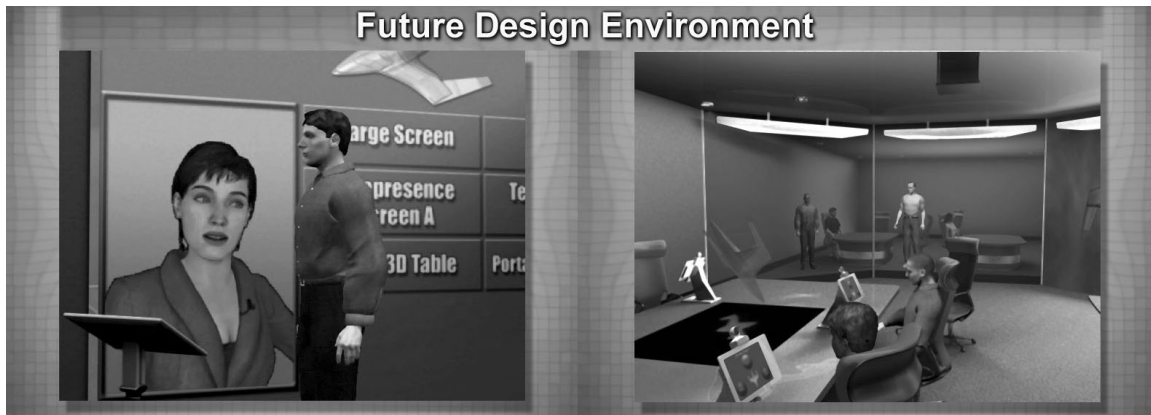


Figure 25

OBJECTIVES AND FORMAT OF WORKSHOP

The objectives of the workshop are to (Figure 26): a) provide broad overviews of the diverse activities related to new computing paradigms, including grid computing, pervasive computing, high-productivity computing, and the IBM-led autonomic computing; and b) identify future directions for research that have high potential for future aerospace workforce environments. The format included twenty half-hour presentations in nine sessions, and three exhibits.

- *Objectives:*
 - Overview of diverse activities related to emerging/new computing paradigms
 - Identify future directions for research for future aerospace workforce environments
- *Format:*
 - 20 presentations; 9 sessions
 - 3 exhibits
- *Proceedings:*
 - NASA Conference Proceeding

Figure 26

INFORMATION ON EMERGING / NOVEL COMPUTING PARADIGMS AND FUTURE COMPUTING ENVIRONMENTS

A short list of books, monographs, conference proceedings, survey papers and websites on emerging/novel computing paradigms and future computing environment is given subsequently.

Books, Monographs and Conference Proceedings:

- [1] Fogg, B.J., *"Persuasive Technology – Using Computers to Change What We Think and Do"*, Morgan Kaufmann Publishers, 2003.
- [2] Grigoras, D., Nicolau, A., Torsel, B. and Folliot, B., (editors), *"Advanced Environments, Tools, and Applications for Cluster Computing"*, NATO Advanced Research Workshop, Iwcc 2001, Mangalis, Romania, September 2001: Revised Papers, Springer-Verlag New York, Inc., May 2002.
- [3] Carroll, J.M. (editor), *"Human-Computer Interaction in the New Millennium"*, ACM Press, New York, 2002.
- [4] Luryi, S., Xu, J., and Zaslavsky, A., *"Future Trends in Microelectronics: The Road Ahead"*, John Wiley and Sons, Inc., 2002.
- [5] Calude, C.S., and Paun, G., *"Computing with Cells and Atoms: An Introduction to Quantum, DNA and Membrane Computing"*, Taylor & Francis, 2001.
- [6] Barkai, D., *"Peer-to-Peer Computing – Technologies for Sharing and Collaborating on the Net"*, Intel Press, 2001.
- [7] Nielsen, M.A. and Chuang, I.L., *"Quantum Computation and Quantum Information"*, Cambridge University Press, 2001.
- [8] Hirvensalo, M., *"Quantum Computing"*, Springer-Verlag TELOS, 2001.
- [9] Kim, D., and Hariri, S., *"Virtual Computing: Concept, Design, and Evaluation"*, Kluwer Academic Publishers, August 2001.
- [10] Greenia, M.W. *"History of Computing: An Encyclopedia of the People and Machines that Made Computer History"*. CD-ROM, Lexikon Services, 2000.
- [11] Ceruzzi, P.E., *"A History of Modern Computing"*, Massachusetts Institute of Technology, 1999.
- [12] Foster, I., and Kesselman, C., (editors), *"The Grid: Blueprint for a New Computing Infrastructure"*, Morgan Kaufmann Publishers, Inc., San Francisco, CA, 1999.
- [13] Denning, P.J., (editor), *"Talking Back to the Machine: Computers and Human Aspiration"*, Springer-Verlag, New York, 1999.
- [14] Kurzweil, R., *"The Age of Spiritual Machines: When Computers Exceed Human Intelligence"*, Penguin Putnam, Inc., New York, 1999.
- [15] Williams, C.P., and Clearwater, S.H., *"Explorations in Quantum Computing"*, Springer-Verlag, New York, 1998.
- [16] Paun, G., Rozenberg, G., Salomaa, A., and Brauer, W., *"DNA Computing: New Computing Paradigms (Texts in Theoretical Computer Science)"*, Springer Verlag, 1998.

- [17] Maybury, M.T., and Wahlster, W., (editors), “*Readings in Intelligent User Interfaces*”, Morgan Kaufmann Publishers, Inc., 1998.
- [18] Denning, P.J., and Metcalfe, R.M., (editors), “*Beyond Calculation: The Next Fifty Years of Computing*”, Springer-Verlag, New York, 1997.
- [19] Williams, M.R., “*A History of Computing Technology*”, IEEE Computer Society; 2nd edition, 1997.
- [20] Shurkin, J., “*Engines of the Mind*”, W.W. Norton & Company, New York. 1996.
- [21] Prince, B., “*High Performance Memories: New Architecture Drams and Srams – Evolution and Function*”, Wiley, John & Sons, Incorporated, 1996.
- [22] Wherrett, B.S., and Chavel, P. (editors), “*Optical Computing*”, Proceedings of the International Conference, Heriot-Watt University, Edinburgh, U.K., August 22 – 25, 1994, Iop Publishers, March 1995.
- [23] AcAulay, A.D., “*Optical Computer Architectures: The Application of Optical Concepts t Next Generation Computers*”, Wiley, John & Sons, Inc., 1991.

Special Issues of Journals:

- [1] Vertegaal, R. “*Attentive User Interfaces*” Editorial, Special Issue on Attentive User Interfaces, Communications of ACM 46(3), March 2003.
- [2] “*Limits of Computation*”, Special issue of Computing in Science and Engineering, May/June 2002.
- [3] “The Future of Computing – Beyond Silicon”, Special Issue of Technology Review, MIT’s Magazine of Innovation, May/June 2000.
- [4] Caulfield, H.J., “*Perspectives in Optical Computing*”, IEEE Computer, Vol. 31, No. 2, February 1998, pp. 22-25.
- [5] “*50 Years of Computing*”, Special issue of IEEE Computer, Vol. 29, No. 10, October 1996.

Survey Papers and Articles:

- [1] Compton, K., and Hauck, S., “*Reconfigurable Computing: A Survey of Systems and Software*”, ACM Computing Surveys, Vol. 34, No. 2, June 2002, pp. 171-210.
- [2] Waldrop, M.M., “*Grid Computing*”, Technology Review, May 2002, pp. 31-37.
- [3] Noor, A.K., “*New Computing Systems and Future High-Performance Computing Environment and Their Impact on Structural Analysis and Design*”, Computers and Structures, Vol. 64, Nos. 1-4, July-August 1997, pp. 1-30.
- [4] Weiser, M., “The Computer of the 21st Century”, Scientific American, Vol. 265, No. 3, September 1991, pp. 66-75.

Websites:

1. MIT Project Oxygen – Pervasive Human-Centered Computing
<http://oxygen.lcs.mit.edu>

2. Autonomic Computing – Creating Self-Managing Computing Systems
<http://www-3.ibm.com/autonomic>
3. Pervasive Computing – Anywhere. Anytime. On Demand.
<http://www.darpa.mil/ipto/research/hpcs>
4. The Globus Project
<http://www.globus.org>
5. Quantum computation: a tutorial
<http://www.sees.bangor.ac.uk/~schmuel/comp/compt.html>
6. Stanford University, U.C. Berkeley, MIT, and IBM Quantum Computation Project
<http://divine.stanford.edu>
7. DNA Computers
http://members.aol.com/ibrandt/dna_computer.html
8. Publications on DNA based Computers
<http://crypto.stanford.edu/~dabo/biocomp.html>
9. European Molecular Computing Consortium (EMCC)
<http://openit.disco.unimib.it/emcc>

AUTONOMIC COMPUTING – AN OVERVIEW

Ric Telford
Autonomic Computing/IBM Corporation
Research Triangle Park, NC

Reprinted with the permission of IBM Corporation.

AUTONOMIC COMPUTING – AN OVERVIEW

This presentation covers the autonomic computing vision and initiative. In 2001, Paul Horn, IBM's Senior VP of Research, issued a "grand challenge" to the computing industry to develop more self-managing systems. Since that time, IBM has been working internally and externally to advance the state of its products and the industry to autonomic computing capabilities. IBM has a corporate initiative to drive autonomic computing. Ric Telford is the Director of Architecture and Technology for this initiative.

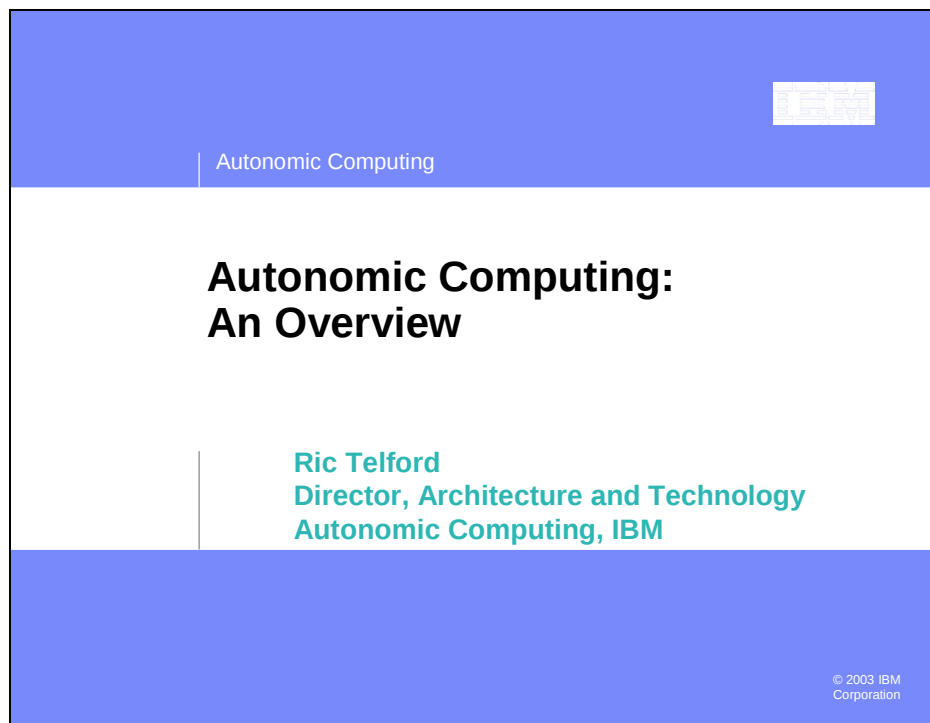


Figure 1

THE E-BUSINESS EVOLUTION

IBM introduced the concept of "e-business on demand" last year to describe the next evolution of e-business. We are now at a stage where the need for e-business on demand is becoming the priority for companies. e-business on demand allows for a more loosely-coupled, service-oriented approach to e-business infrastructures. It allows for systems to combine and separate as required in real-time, to address the business problems at hand. It allows for the leveraging of systems not only within a data center, but across multiple data centers, other businesses, and service providers. This is the vision of computing - being able to construct systems easier, and being able to run these systems with minimal downtime and minimal human intervention. IBM realized a while back that we need some new, fundamental technologies to make this vision of dynamic e-business real. Specifically, the ability for IT systems to manage themselves - react to problems, re-configure based on load, guard against attacks and continually optimize based on set policies. This is autonomic computing.

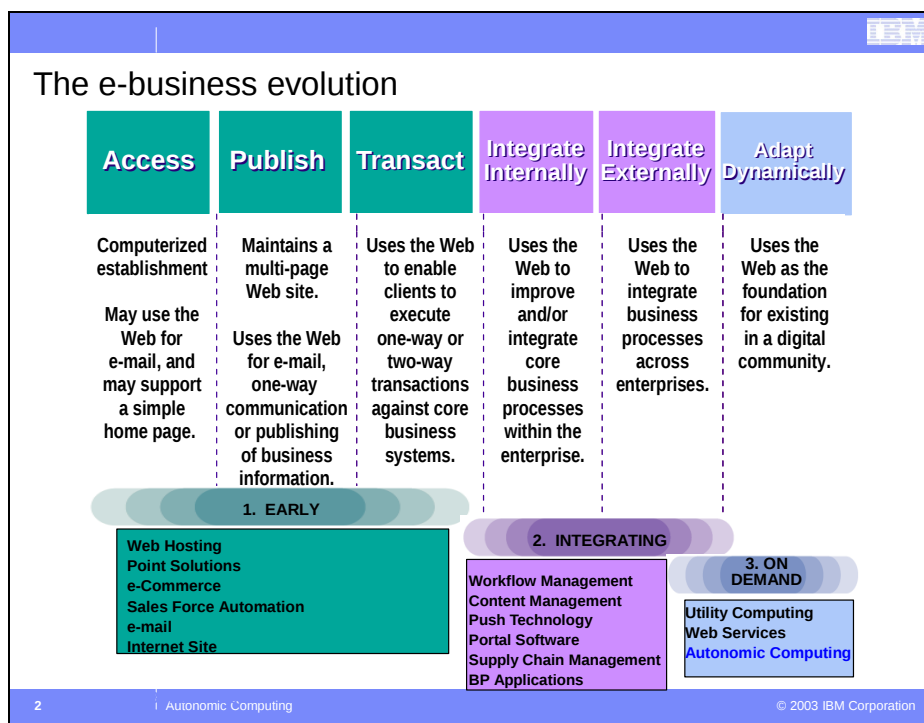


Figure 2

AUTONOMIC COMPUTING AS PART OF E-BUSINESS ON DEMAND

An on-demand business has some key qualities - the ability to respond to requests/demands in real-time. The ability to have variable cost structures in IT vs. the static (fixed) cost structures of today. The ability to have the time-consuming tasks of managing an IT system be more self-managed, allowing for IT professionals to work on that which is core and differentiating for the business. Finally, the ability for a system to be resilient and highly available.

To build such a system, you need an operating environment with some key capabilities - the ability to integrate across systems, open standards-based, a virtualized infrastructure and of course, autonomic computing capabilities.

Autonomic Computing as part of e-business On Demand

@business on demand
The *on demand* era is upon us.
Are you ready?

On Demand Business

- Responsive in real-time
- Variable cost structures
- Focused on what's core and differentiating
- Resilient around the world, around the clock

On Demand Operating Environment

- Integrated
- Open
- Virtualized
- **Autonomic**

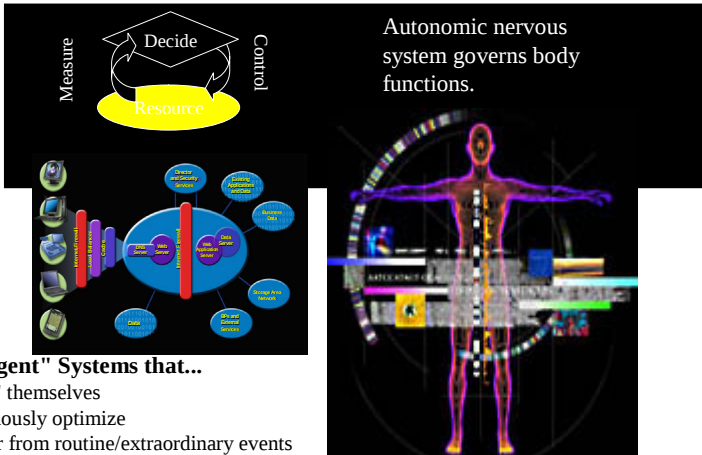
3 | Autonomic Computing © 2003 IBM Corporation

Figure 3

WHY AUTONOMIC COMPUTING?

The term Autonomic Computing comes from the human autonomic nervous system. The autonomic system “self-manages” the body (heart rate, breathing, etc). Computing systems need to be able to do the same.

Why Autonomic Computing?



The diagram on the left illustrates the autonomic nervous system cycle, showing a loop between 'Measure' and 'Control' leading to 'Decide', which then leads to 'Resource'. The diagram on the right shows a human figure with internal organs highlighted, representing the autonomic nervous system governing body functions.

Autonomic nervous system governs body functions.

"Intelligent" Systems that...

- "Know" themselves
- Continuously optimize
- Recover from routine/extraordinary events
- Anticipate and adapt to user needs
- Protect against attacks threatening the system
- Understand the external environment in which they operate
- Support heterogeneous environments via open standards
- Configure and re-configure under varying and unpredictable conditions

4 | Autonomic Computing © 2003 IBM Corporation

Figure 4

COMPLEXITY

The need for autonomic computing is required today due to the ever growing complexity of systems. Keeping systems properly configured, optimized and running is a very labor intensive task. Autonomic computing is focused at reducing this complexity in the IT infrastructure.

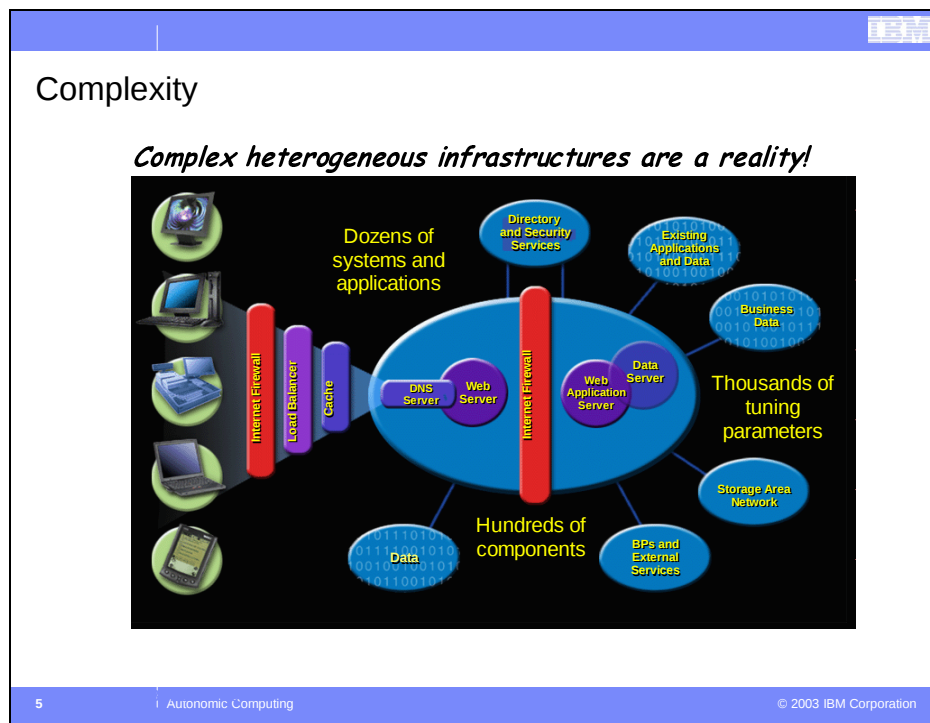
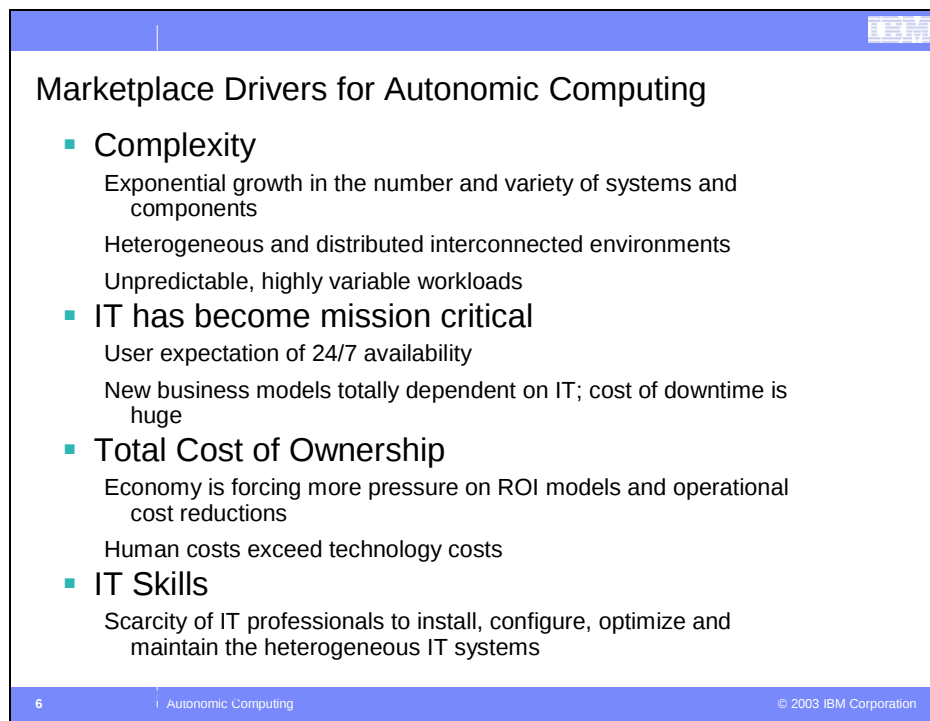


Figure 5

MARKETPLACE DRIVERS FOR AUTONOMIC COMPUTING

Although complexity is the most significant driver for autonomic computing, there are others as well.



Marketplace Drivers for Autonomic Computing

- **Complexity**
 - Exponential growth in the number and variety of systems and components
 - Heterogeneous and distributed interconnected environments
 - Unpredictable, highly variable workloads
- **IT has become mission critical**
 - User expectation of 24/7 availability
 - New business models totally dependent on IT; cost of downtime is huge
- **Total Cost of Ownership**
 - Economy is forcing more pressure on ROI models and operational cost reductions
 - Human costs exceed technology costs
- **IT Skills**
 - Scarcity of IT professionals to install, configure, optimize and maintain the heterogeneous IT systems

6 | Autonomic Computing | © 2003 IBM Corporation

Figure 6

WHAT IS AUTONOMIC COMPUTING?

We talk about self-managing systems in four areas: 1) self configuring, or the ability to understand the environment and configure accordingly, 2) self-healing, which is the ability for systems to determine problems and workaround or fix the problem, 3) self-optimizing, the ability to re-configure based on changing conditions and 4) self-protecting which is the ability to guard against external threats.

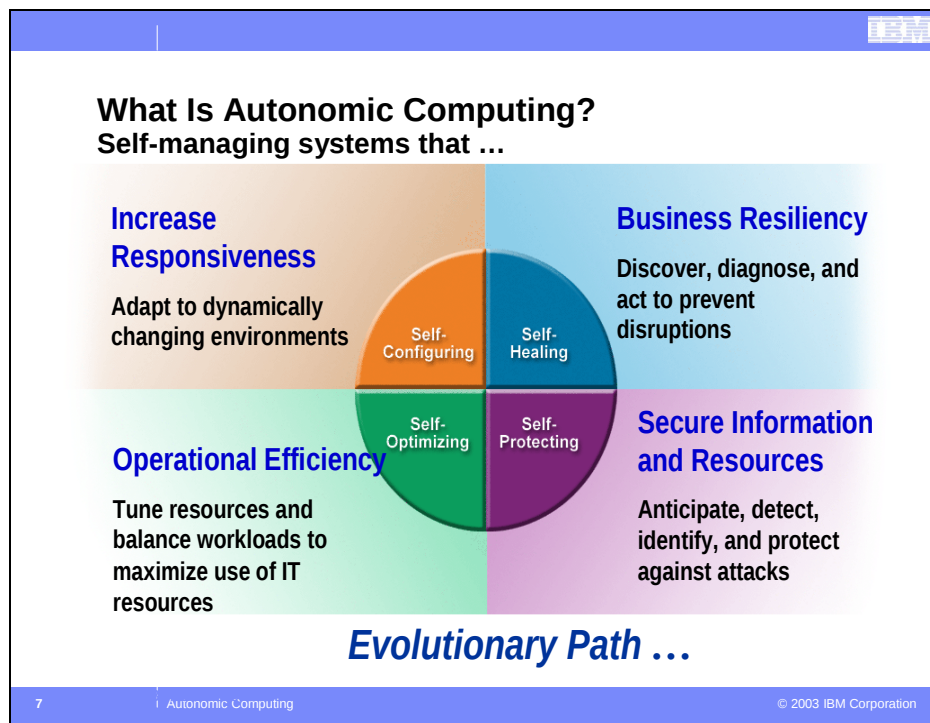


Figure 7

EVOLUTION; NOT REVOLUTION

A fully autonomic system is something you evolve to from where you are today. This chart lays out the progression from a manual IT infrastructure to a fully autonomic one.

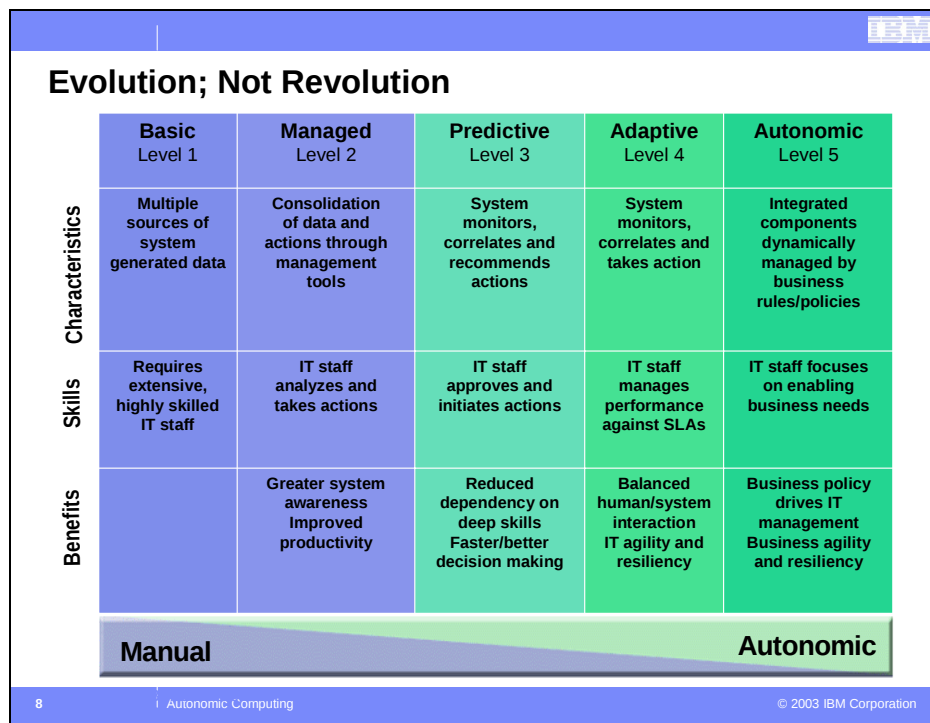


Figure 8

AN ANALOGY – CAMERA “AUTONOTMICS”

There is an analogy here to the camera industry. Over time, cameras evolved from a very manual set of functions to a highly automated set. One point of note here - even though the automated capabilities exist in the camera industry, it is still possible to configure the camera manually. This will be true in autonomic computing systems as well.

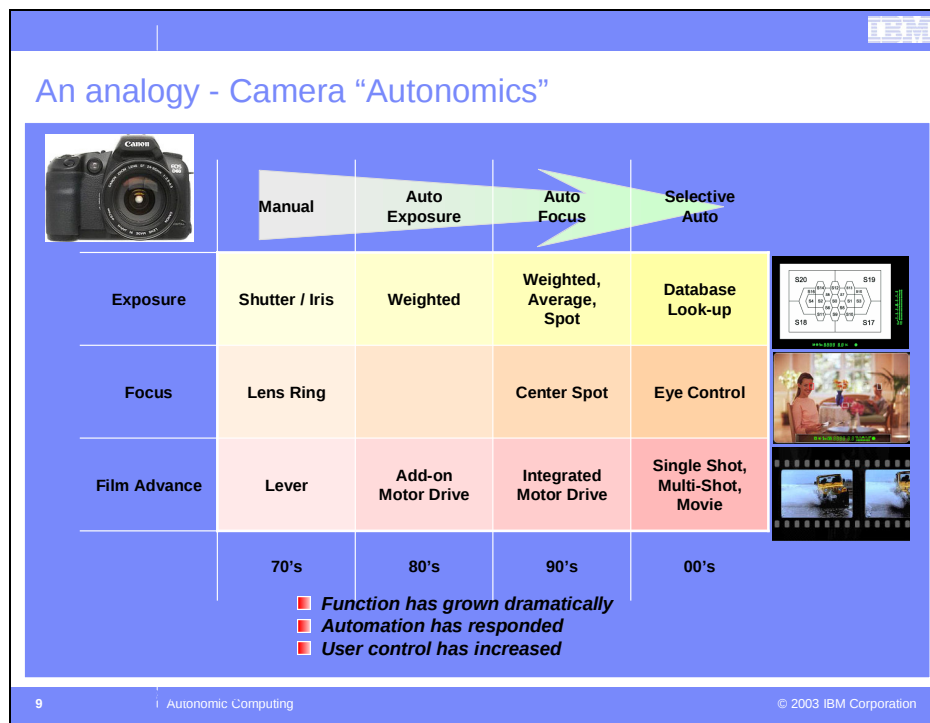


Figure 9

AUTONOMIC COMPUTING REFERENCE MODEL

This diagram is a model for what an autonomic manager looks like. It requires elements of a system (database, server, storage, applications) to expose a set of “sensors” (state information on the element) and “effectors” (interfaces for tuning, configuring, changing state, etc). Given any set of sensors and effectors, an autonomic manager can be built which monitors the sensors, analyzes the data, compares the existing state to the desired state (rendered as “knowledge”) and then set a plan and execution for change.

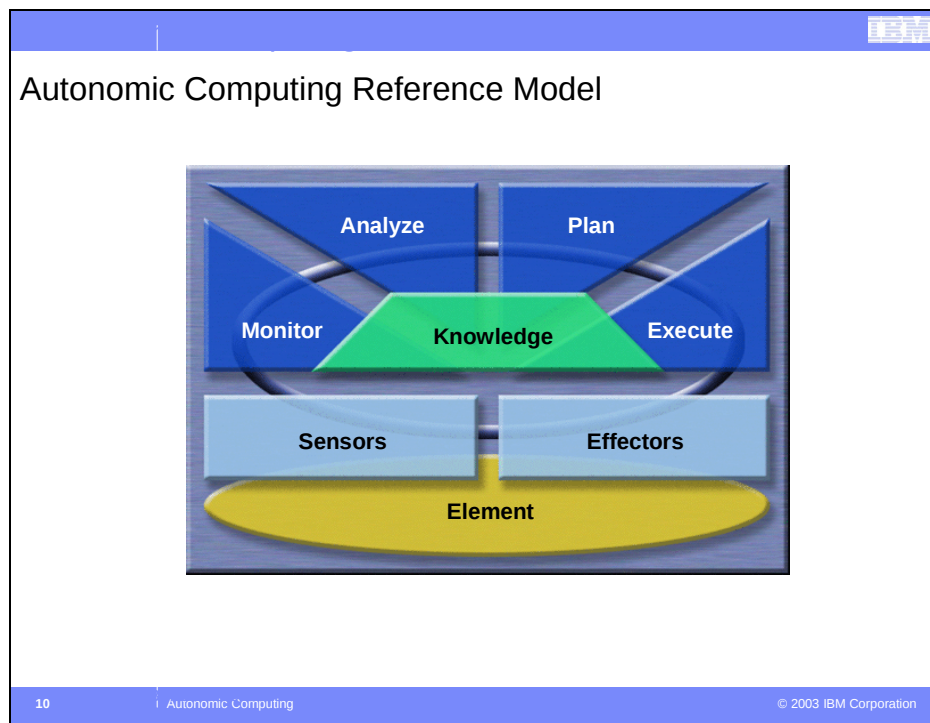


Figure 10

MULTIPLE CONTEXTS FOR AUTONOMIC BEHAVIOR

Autonomic managers can exist at many layers in the system. The challenge is to coordinate the behaviors of the AC systems.

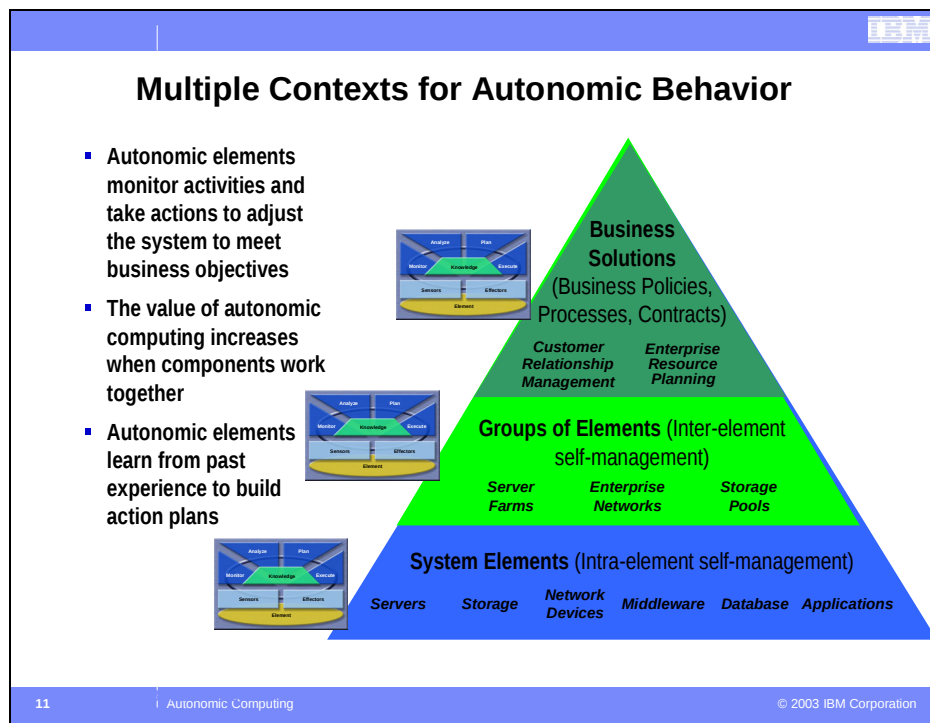


Figure 11

SELF – CONFIGURING EXAMPLE

This is a very good example of a self-configuring system. The Configuration Advisor in DB2 can self-configure a database system. Often the results of the Configuration Advisor are as good as, if not better, than a human database administrator.

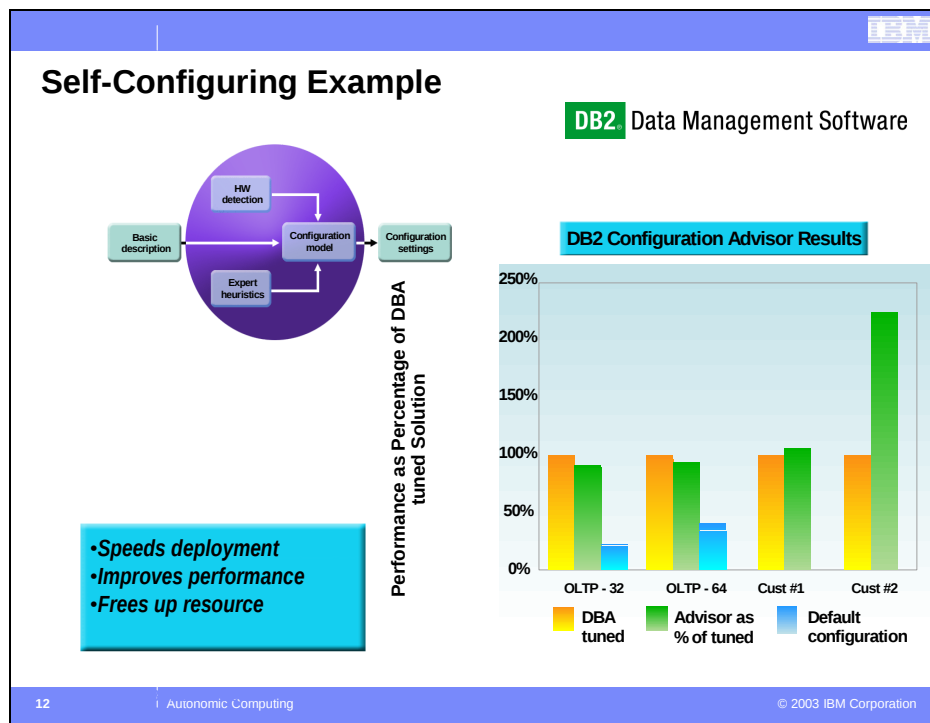


Figure 12

SELF-PROTECTING EXAMPLE

Tivoli's Risk Manager is an example of a self-protecting system. By monitoring and correlating data from across the infrastructure, Risk Manager can determine if there are external threats to the system.

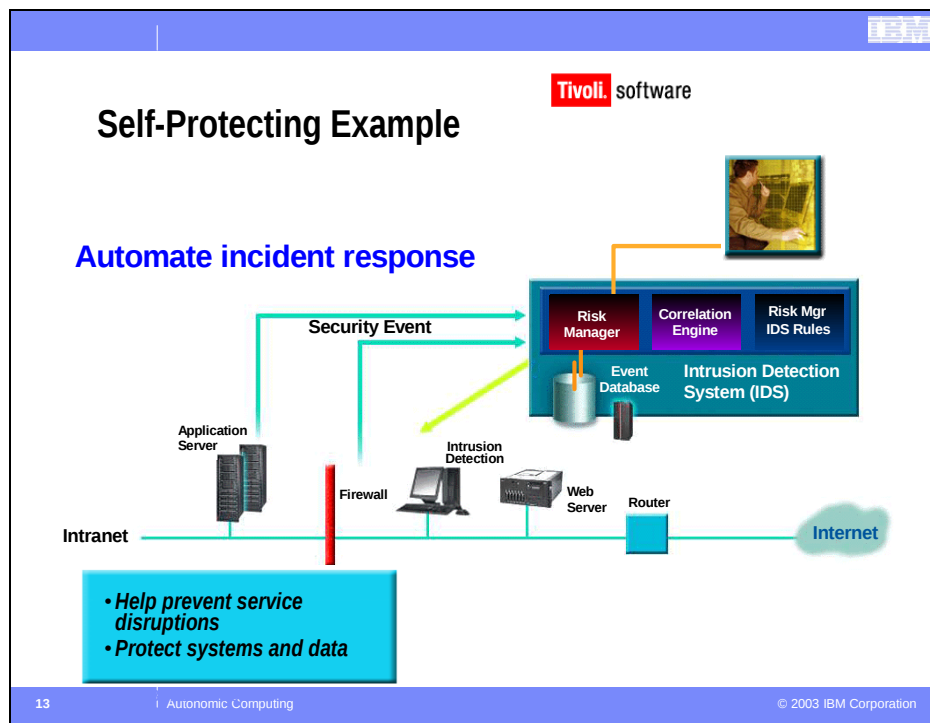


Figure 13

AUTONOMIC COMPUTING WILL IMPACT IT PROCESSES

Finally, it is important to note that autonomic computing will greatly help in the automation of processes in an IT infrastructure.

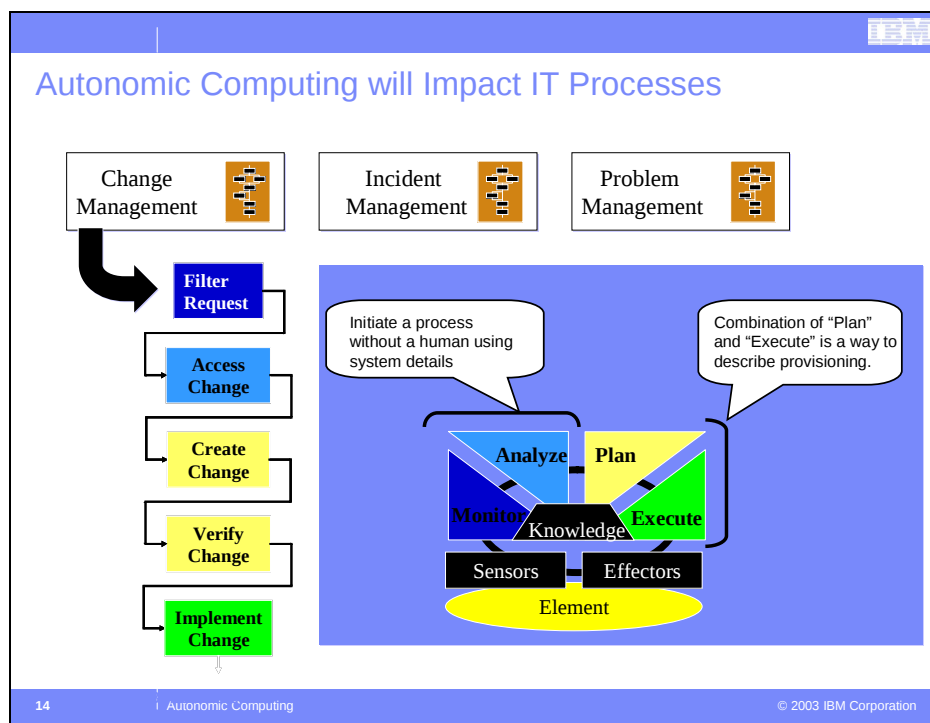


Figure 14

**Towards Autonomic Computational Science & Engineering
(Autonomic Computing: Application Perspective)**

Manish Parashar
The Applied Software Systems Laboratory
ECE/CAIP, Rutgers University

THE CURRENT TEAM

The AutoMate team is composed of faculty and graduate and undergraduate students at The Applied Software Systems Laboratory, Department of Electrical and Computer Engineering and Center of Advanced Information Processing (CAIP), Rutgers, The State University of New Jersey. The team is organized as the Autonomic Computing Research Group and the Autonomic Applications Research Group. This research builds on our collaborations with application scientists, engineers and computer and computational scientists at California Institute of Technology, University of Texas at Austin, University of Arizona, Ohio State University, and University of Maryland.

The Current Team

- TASSL Rutgers University
 - Autonomic Computing Research Group
 - Viraj Bhat
 - Manish Agarwal
 - Hua Liu (Maria)
 - Zhen Li (Jenny)
 - Manish Mahajan
 - Vincent Matossian
 - Venkatesh Putty
 - Cristina Schmidt
 - Guangsen Zhang
 - Autonomic Applications Research Group
 - Sumir Chandra
 - Xiaolin Li
 - Taher Saif
 - Li Zhang
 - Hailan Zhu
- CS Collaborators
 - HPDC, University of Arizona
 - Salim Hariri
 - Biomedical Informatics, The Ohio State University
 - Tahsin Kurc, Joel Saltz
 - CS, University of Maryland
 - Alan Sussman, Christian Hansen
- Applications Collaborators
 - CSM, University of Texas at Austin
 - Malgorzata Peszynska, Mary Wheeler
 - IG, University of Texas at Austin
 - Mrinal Sen, Paul Stoffa
 - ASCI/CACR, Caltech
 - Michael Aivazis, Julian Cummings, Dan Meiron
 - CRL, Sandia National Laboratory, Livermore
 - Jaideep Ray, Johan Steensland



Figure 1

OVERVIEW OF THE TALK

This talk motivates and introduces autonomic computational science and engineering, and presents the AutoMate framework for enabling autonomic applications on Grid. It describes the AutoMate architecture and briefly presents each of its components. These include the ACCORD autonomic component framework, the RUDDER decentralized deductive engine, the SESAME context-sensitive dynamic access management framework, the Pawn peer-to-peer messaging substrate, and the SQUID decentralized discovery service. Finally, it describes two applications of autonomic computing to science and engineering – autonomic runtime management framework for adaptive applications (V-Grid) and autonomic interactions for oil reservoir optimization.

Outline

- Autonomic computational science and engineering
- AutoMate: A framework of enabling autonomic applications
 - *ACCORD: Autonomic component framework*
 - *RUDDER: Decentralized deductive engine*
 - *SESAME: Context sensitive dynamics access management*
 - *Pawn: Peer-to-Peer messaging infrastructure*
 - *SQUID: Decentralized discovery service*
- Application Scenarios
 - *V-Grid autonomic runtime for adaptive applications*
 - *reactive/proactive partitioning, load-balancing, scheduling, performance management*
 - *Autonomic interactions oil reservoir optimization*
- Conclusions and current status



Figure 2

COMPUTATION MODELING OF PHYSICAL PHENOMENA

Realistic, physically accurate simulations of complex physical phenomena that symbiotically and opportunistically combine computations, experiments, observations, and real-time data have the potential for providing dramatic insights into complex systems such as interacting black holes and neutron stars, formations of galaxies, subsurface flows in oil reservoirs and aquifers, and dynamic response of materials to detonations. However, the phenomena being modeled by these applications are inherently large-scale, dynamic and heterogeneous (in time, space, and state). Furthermore, the applications are extremely large with unprecedented resource requirements, and are composed of a large numbers of software components with very dynamic compositions and interactions between these components.

Computational Modeling of Physical Phenomenon

- Realistic, physically accurate computational modeling
 - Large computation requirements
 - e.g. simulation of the core-collapse of supernovae in 3D with reasonable resolution (500^3) would require ~ 10 - 20 teraflops for 1.5 months (i.e. ~ 100 Million CPUs!) and about 200 terabytes of storage
 - e.g. turbulent flow simulations using active flow control in aerospace and biomedical engineering requires $5000 \times 1000 \times 500 = 2.5 \cdot 10^9$ points and approximately 107 time steps, i.e. with 1GFlop processors requires a runtime of $\sim 7 \cdot 10^6$ CPU hours, or about one month on 10,000 CPUs! (with perfect speedup). Also with 700B/pt the memory requirement is ~ 1.75 TB of run time memory and ~ 800 TB of storage.
 - Coupled, multiphase, heterogeneous, dynamic
 - multi-physics, multi-model, multi-resolution,
 - Complex interactions
 - application – application, application – resource, application – data, application – user, ...
 - Software/systems engineering/programmability
 - volume and complexity of code, community of developers, ...
 - scores of models, hundreds of components, millions of lines of code, ...



Figure 3

COMPUTATION MODELING AND THE GRID

The emergence of computational Grids and the potential for seamless aggregation, integration and interactions has made it possible to conceive the realistic, scientific and engineering simulations of complex physical phenomena described in the previous slide. However, the Grid infrastructure is also heterogeneous and dynamic, globally aggregating large numbers of independent computing and communication resources, data stores and sensor networks. The combination of the two (large, complex, heterogeneous and dynamic applications and Grids) results in application development, configuration and management complexities that break current paradigms based on passive components and static compositions. Clearly, there is a need for a fundamental change in how these applications are formulated, composed and managed so that their heterogeneity and dynamics can match and exploit the heterogeneous and dynamic nature of the Grid. In fact, we have reached a level of complexity, heterogeneity, and dynamism for which our programming environments and infrastructure are becoming unmanageable brittle and insecure. This has led researchers to consider alternative programming paradigms and management techniques that are based on strategies used by biological systems to deal with complexity, heterogeneity and uncertainty. The approach is referred to as autonomic computing. An autonomic computing system is one that has the capabilities of being self-defining, self-healing, self-configuring, self-optimizing, self-protecting, context aware, and open.

Computational Modeling and the Grid

- The Computational Grid
 - Potential for aggregating resources
 - computational requirements
 - Potential for seamless interactions
 - new applications formulations
- Developing application to utilize and exploit the Grid remains a significant challenge
 - The problem: a level of complexity, heterogeneity, and dynamism for which our programming environments and infrastructure are becoming unmanageable, brittle and insecure
 - System size, heterogeneity, dynamics, reliability, availability, usability
 - Currently typically proof-of-concept demos by "hero programmers"
 - Requires fundamental changes in how applications are formulated, composed and managed
 - Breaks current paradigms based on passive components and static compositions
 - autonomic components and their dynamic composition, opportunistic interactions, virtual runtime, ...
 - Resonance - heterogeneity and dynamics must match and exploit the heterogeneous and dynamic nature of the Grid
- Autonomic, adaptive, interactive simulations and the Grid offer the potential for such simulations
 - Autonomic: context aware, self configuring, self adapting, self optimizing, self healing,...
 - Adaptive: resolution, algorithms, execution, scheduling, ...
 - Interactive: peer interactions between computational objects and users, data, resources, ...



Figure 4

AUTOMATE

The overall objective of the AutoMate project is to investigate key technologies to enable the development of autonomic Grid applications that are context aware and are capable of self-configuring, self-composing, self-optimizing and self-adapting. Specifically, it will investigate the definition of autonomic components, the development of autonomic applications as dynamic composition of autonomic components, and the design of key enhancements to existing Grid middleware and runtime services to support these applications.

Definition of Autonomic Components: The definition of programming abstractions and supporting infrastructure that will enable the definition of autonomic components. In addition to the interfaces exported by traditional components, autonomic components provide enhanced profiles or contracts that encapsulate their functional, operational, and control aspects. These aspects export information and policies about their behavior, resource requirements, performance, interactivity and adaptability to system and application dynamics. Furthermore, they encapsulate sensors, actuators, access policies and a policy-engine. Together, aspects, policies, and policy engine allow autonomic components to consistently configure, manage, adapt and optimize their execution.

Dynamic Composition of Autonomic Applications: The development of mechanisms and supporting infrastructure to enable autonomic applications to be dynamically and opportunistically composed from autonomic components. The composition will be based on policies and constraints that are defined, deployed and executed at run time, and will be aware of available Grid resources (systems, services, storage, data) and components, and their current states, requirements, and capabilities.

Autonomic Middleware Services: The design, development, and deployment of key services on top of the Grid middleware infrastructure to support autonomic applications. One of the key requirements for autonomic behavior and dynamic compositions is the ability of the components, applications and resources (systems, services, storage, data) to interact as peers. Furthermore the components should be able to sense their environment. In this project, we extend the Grid middleware with (1) a peer-to-peer messaging substrate, (2) context aware services, and (3) peer-to-peer deductive engines for composition, configuration and management of autonomic applications. An active peer-to-peer control network will combine sensors, actuators and rules to configure and tune components and their execution environment at runtime and to satisfy requirements and performance and quality of service constraints.

AutoMate: Enabling Autonomic Applications

- Objective:
 - Investigate key technologies to enable the development of autonomic Grid applications that are context aware and are capable of self-configuring, self-composing, self-optimizing and self-adapting.
- Overview:
 - Definition of Autonomic Components:
 - definition of programming abstractions and supporting infrastructure that will enable the definition of autonomic components
 - autonomic components provide enhanced profiles or contracts that encapsulate their functional, operational, and control aspects
 - Dynamic Composition of Autonomic Applications:
 - mechanisms and supporting infrastructure to enable autonomic applications to be dynamically and opportunistically composed from autonomic components
 - compositions will be based on policies and constraints that are defined, deployed and executed at run time, and will be aware of available Grid resources (systems, services, storage, data) and components, and their current states, requirements, and capabilities
 - Autonomic Middleware Services:
 - design, development, and deployment of key services on top of the Grid middleware infrastructure to support autonomic applications
 - a key requirements for autonomic behavior and dynamic compositions is the ability of the components, applications and resources (systems, services, storage, data) to interact as peers



Figure 5

AutoMate: Architecture

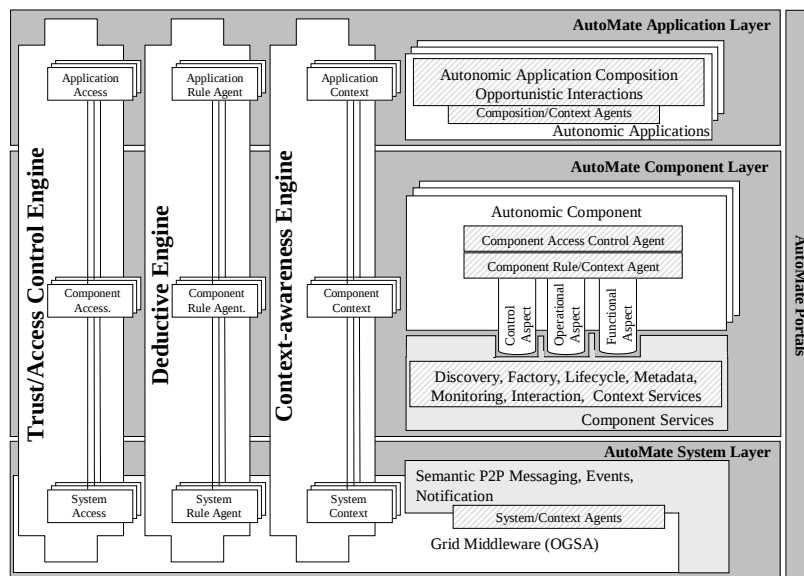


Figure 6

AUTOMATE ARCHITECTURE

AutoMate builds on the emerging Grid infrastructure and extends the Open Grid Service Architecture (OGSA). AutoMate is composed of the following components:

AutoMate System Layer: The AutoMate system layer builds on the Grid middleware and OGSA and extends core Grid services (security, information and resource management, data management) to support autonomic behavior. Furthermore, this layer provides specialized services such as peer-to-peer semantic messaging, events and notification.

AutoMate Component Layer: The AutoMate component layer addresses the definition, execution and runtime management of autonomic components. It consists of AutoMate components that are capable of self configuration, adaptation and optimization, and supporting services such as discovery, factory, lifecycle, context, etc. (which builds on core OGSA services).

AutoMate Application Layer: The AutoMate application layer builds on the component and system layers to support the autonomic composition and dynamic (opportunistic) interactions between components.

AutoMate Engines: AutoMate engines are decentralized (peer-to-peer) networks of agents in the system. The context-awareness engine is composed of context agents and services and provides context information at different levels to trigger autonomic behaviors. The deductive engine is composed of rule agents which are part of the applications, components, services and resources, and provides the collective decision making capability to enable autonomic behavior. Finally, the trust and access control engine is composed of access control agents and provides dynamic context-aware control to all interactions in the system.

In addition to these layers, AutoMate portals provide users with secure, pervasive (and collaborative) access to the different entities. Using these portals users can access resource, monitor, interact with, and steer components, compose and deploy applications, configure and deploy rules, etc. AutoMate leverages the experiences and technologies developed as part of the Discover/DIOS computational collaboratory project (<http://www.discoverportal.org>). The different components are described in the following sections.

AUTOMATE ARCHITECTURE

Key components of AutoMate include:

- ACCORD (Autonomic Components, Compositions and Coordination) component framework that enables the definition of autonomic components, their autonomic compositions and opportunistic interactions.
- RUDDER (Rule Definition Deployment and Execution Service) decentralized deductive engine.
- SESAME (Scalable Environment Sensitive Access Management Engine) dynamic access control engine.
- Pawn decentralized (P2P) messaging substrate.
- SQUID flexible) information discovery service.

These components are introduced in the following slides.

AutoMate: Components

- **ACCORD: Autonomic application framework**
- **RUDDER: Decentralized deductive engine**
- **SESAME: Dynamic access control engine**
- **Pawn: P2P messaging substrate**
- **SQUID: P2P discovery service**



Figure 7

ACCORD: AUTONOMIC COMPONENTS

Autonomic components in AutoMate export information and policies about their behavior, resource requirements, performance, interactivity, and adaptability to system and application dynamics. In addition to the functional interfaces exported by traditional components, AutoMate components provide semantically enhanced profiles or contracts that encapsulate their functional, operational, and control aspects. A conceptual overview of an AutoMate component is presented in the figure. The functional aspect specification abstracts component functionality, such as order of interpolation (linear, quadratic, etc.). This functional profile is then used by the compositional engine to select appropriate components based on application requirements. The operational aspect specification abstracts a component's operational behavior, including computational complexity, resource requirements, and performance (scalability). This profile is then used by the configuration and runtime engines to optimize component selection, mapping and adaptation. Finally, the control aspect describes the adaptability of the component and defines sensors/actuators and policies for management, interaction and control.

AutoMate components also encapsulate access policies, rules, a rule agent, and an access agent that allow the components to consistently and securely configure, manage, adapt and optimize their execution based on rules and access policies. The access agent is a part of the AutoMate access control engine and the underlying dynamic access control model, and manages access to the component based on its current context and state. The rule agent is part of RUDDER, the AutoMate deductive engine and manages local rule definition, evaluation and execution at the component level. Rules can be dynamically defined (and changed) in terms of the component's interfaces (based on access policies) and system and environmental parameters. Execution of rules can change the state, context and behavior of a component, and can generate events to trigger other rule agents.

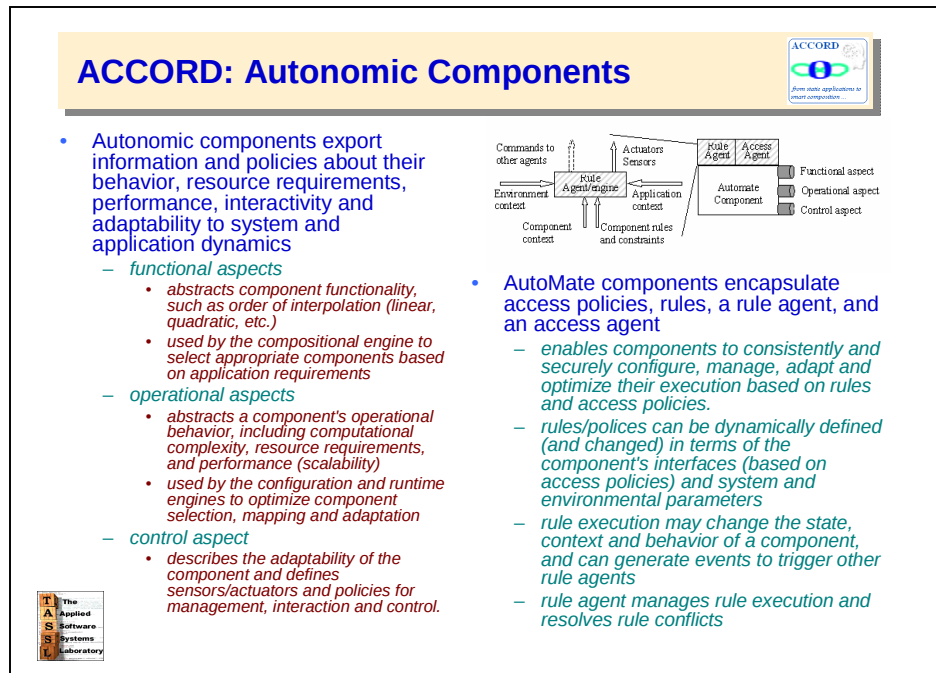


Figure 8

ACCORD: AUTONOMIC COMPOSITIONS

Applications are typically composed with well defined objectives. In case of autonomic applications, however, these objectives can dynamically change based on the state of the application and/or the system. As a result, we need to dynamically select components and compose them at runtime based on current objectives. Together, the profiles, policies, and rules allow autonomous components to consistently and securely manage and optimize their executions. Furthermore, they enable applications to be dynamically composed, configured and adapted. Dynamic application work-flows can be defined to select the most appropriate components based on user/application constraints (highest-performance, lowest cost, reservation, execution time upper bound, best accuracy), on the current applications requirements, to dynamically configure the component's algorithms and behavior based on available resources or system and/or applications state, and to adapt this behavior if necessary.

The AutoMate dynamic composition model may be viewed as transforming a given composition or workflow into a new one by adding or modifying interactions and participating entities. Its primary goal is to enable dynamic (and opportunistic) choreography and interactions of components and services to react to the heterogeneity and dynamics of the application and underlying execution environment to produce the desired user objectives.

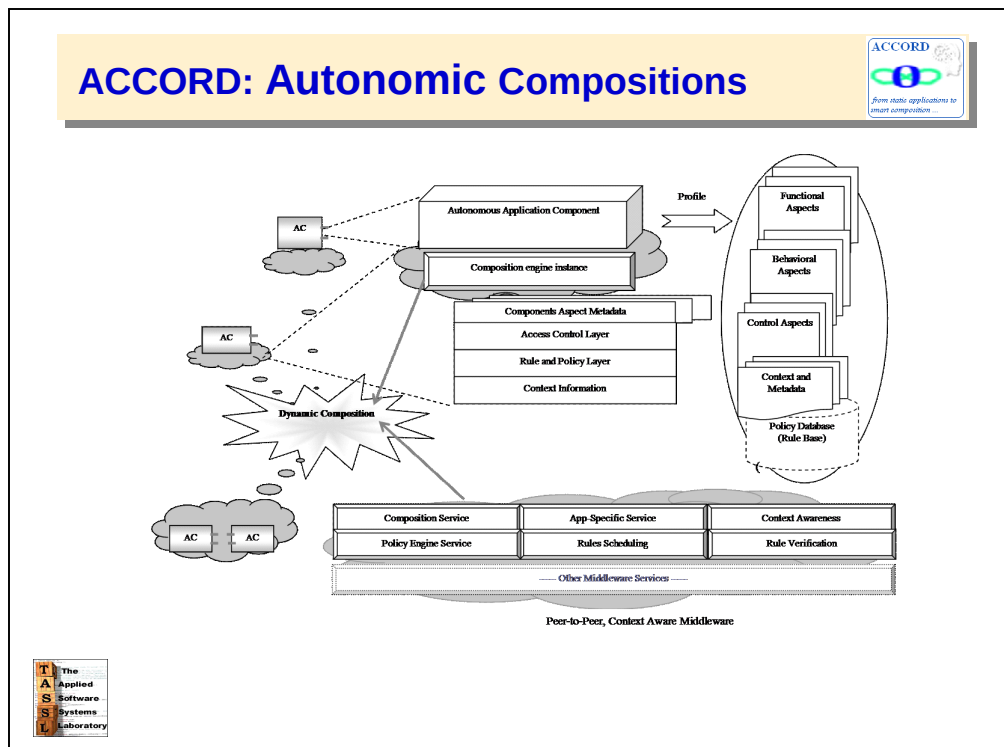


Figure 9

ACCORD: OPPORTUNISTIC INTERACTIONS

Opportunistic interactions are decentralized and based on the satisfaction of locally defined goals and constraints. These interactions are inherently dynamic and ad-hoc and use semantic publisher/subscriber messaging based on proximity, privileges, capabilities, context, interests, and offerings. The goals/constraints are typically long-term and may or may not be satisfied. The interactions do not involve explicit synchronization – the semantics are achieved through feedback and consensus building mechanisms.

ACCORD: Opportunistic Interactions



- Interactions based on local goals and objectives
 - *local goals and objectives are defined as constraints that to be satisfied*
 - *constraints can updated and new constraints can defined at any time*
- Dynamic and ad-hoc
 - *interactions use “semantic messaging” based on proximity, privileges, capabilities, context, interests, offerings, etc.*
- Opportunistic
 - *constraints are long-term and satisfied opportunistically (may not be satisfied)*
- Probabilistic guarantees and soft state
 - *no explicit synchronization*
 - *interaction semantics are achieved using feedback and consensus building*



Figure 10

RUDDER: DEDUCTIVE ENGINE

RUDDER provides the core capabilities for supporting autonomic compositions, adaptations, and optimizations. It is a decentralized deductive engine composed of distributed specialized agents (component rule agents, composition agents, context agents and system agents) that exist at different levels of the system, and represents their collective behavior. It provides mechanisms for dynamically defining, configuring, modifying and deleting rules. Furthermore it defines an XML schema for composing rules and provides mechanisms for deploying and routing rules, decomposing and distributing them to relevant agents, and for coordinating the execution of rules. It also manages conflict resolutions within a single entity and across entities.

The figure presents a schematic overview of RUDDER. It builds on AutoMate and Grid services and the underlying semantic messaging infrastructure. Rules can be dynamically injected into the system and are routed by the messaging substrate to the appropriate agents. Furthermore, the agents may hierarchically decompose a rule and distribute it to peer agents. For example, an application level rule may be decomposed into sub-rules that are assigned to its components. The components rules may be further decomposed into rules for the underlying systems entities.

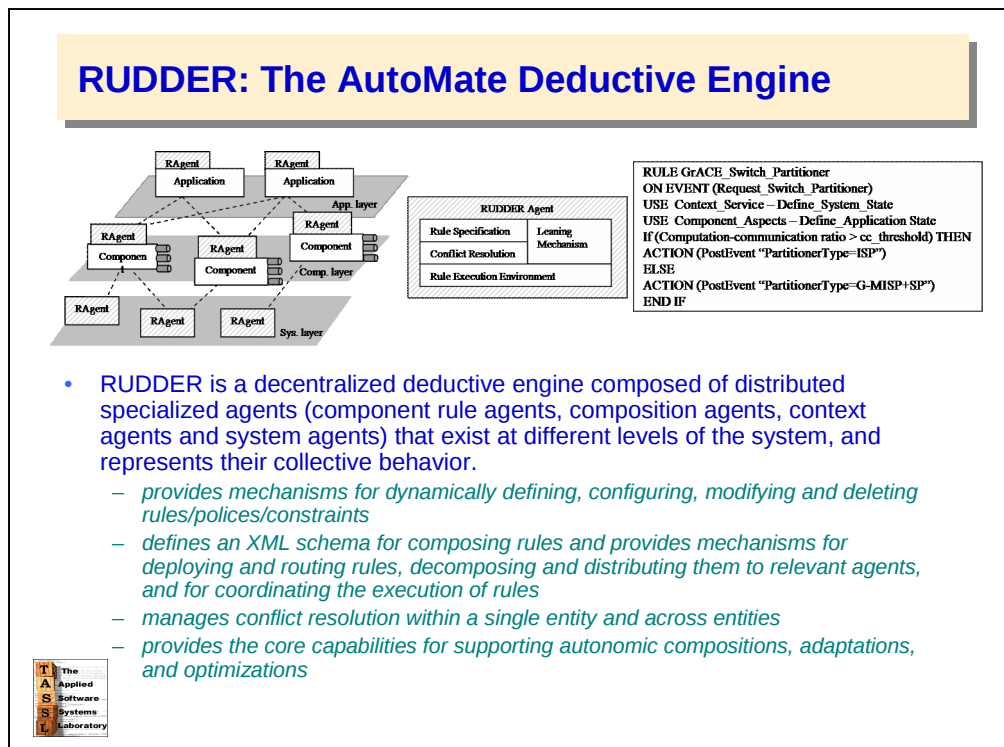


Figure 11

SESAME: CONTEXT SENSITIVE ACCESS MANAGEMENT

A key requirement of autonomic applications is the support for dynamic, seamless and secure interactions between the participating entities, i.e. components, services, application, data, instruments, resources and users. Ensuring interaction security requires a fine grained access control mechanism. Furthermore, in the highly dynamic and heterogeneous Grid environment, the access rights of an entity depend on the entity's privileges, capabilities, context and state. For example, the ability of a user to access a resource or steer a component depends on users' privileges (e.g. owner), current capabilities (e.g. resources available), current context (e.g. secure connection) and the state of the resource or component. The AutoMate Access Control Engine addresses these issues and provides dynamic access control to users, applications, services, components and resources. The engine is composed of access control agents associated with various entities in the system. The underlying dynamic role based access control mechanism extends the RBAC (Role Based Access Control) model to make access control decision based on dynamic context information. The access control engine dynamically adjusts *Role Assignments* and *Permission Assignments*.

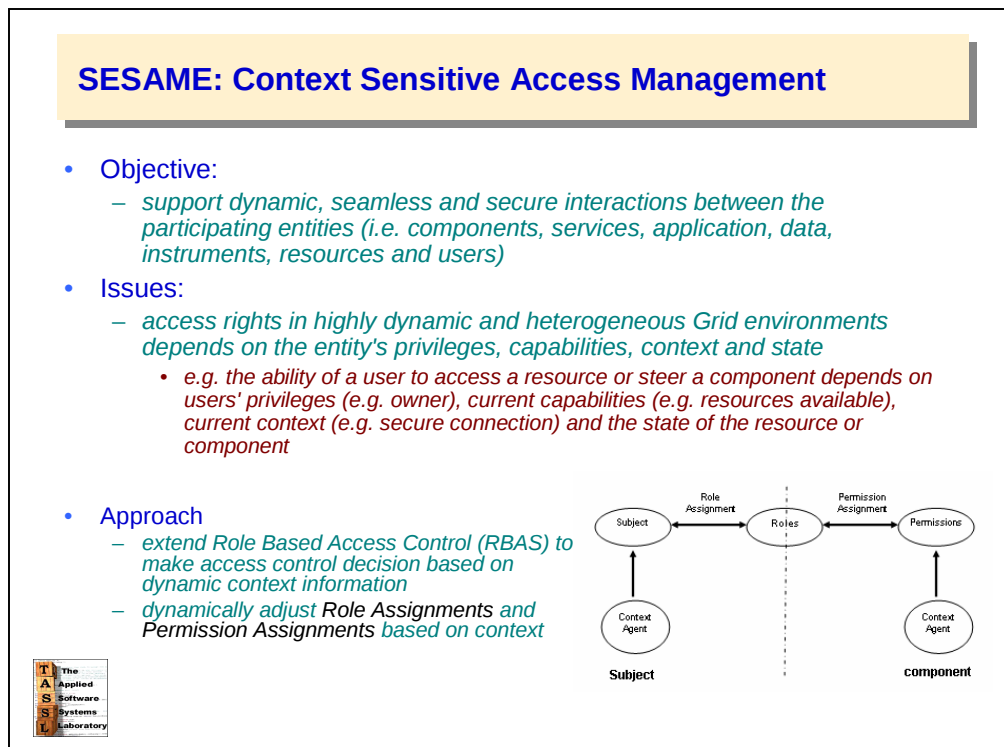


Figure 12

PAWN: P2P MESSAGING

Pawn is a peer-to-peer messaging substrate that builds on project JXTA to support peer-to-peer interactions on the Grid. Pawn provides a stateful and guaranteed messaging to enable key application-level interactions such as synchronous/asynchronous communication, dynamic data injection, and remote procedure calls. It exports these interaction modalities through services at every step of the scientific investigation process, from application deployment, to interactive monitoring and steering, and group collaboration.

A conceptual overview of the Pawn P2P substrate is presented in the figure. Pawn is composed of peers (computing, storage, or user peers), network and interaction services, and mechanisms. These components are layered to represent the requirements stack enabling interactions in a Grid environment. The figure can be read from bottom to top as: “Peers compose messages handled by services through specific interaction modalities”.

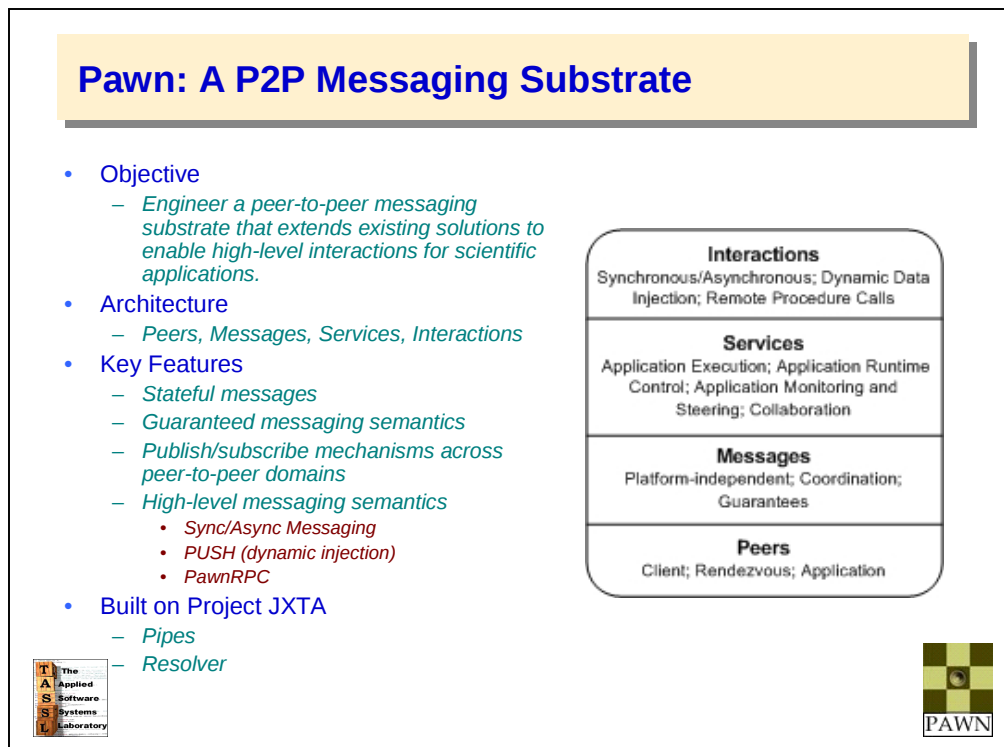


Figure 13

SQUID: DECENTRALIZED Discover

A fundamental problem in large, decentralized, distributed resource sharing environments such as the Grid is the efficient discovery of information, in the absence of global knowledge of naming conventions. For example a document is better described by keywords than by its filename, a computer by a set of attributes such as CPU type, memory, operating system type than by its host name, and a component by its aspects than by its instance name. The heterogeneous nature and large volume of data and resources, their dynamism (e.g. CPU load) and the dynamism of the Grid make the information discovery a challenging problem. An ideal information discovery system has to be efficient, fault-tolerant, self-organizing, has to offer guarantees and support flexible searches (using keywords, wildcards, range queries). Decentralized peer-to-peer (P2P) systems, by their inherent properties (self-organization, fault-tolerance, scalability), provide an attractive solution.

SQUID supports decentralized information discovery in AutoMate. It is a P2P system that supports complex queries containing partial keywords, wildcards, and range queries, and guarantees that all existing data elements that match a query will be found with bounded costs in terms of number of messages and number of nodes involved. The key innovation is a dimension reducing indexing scheme that effectively maps the multidimensional information space to physical peers.

SQUID: A Decentralized Discovery Service

- Overview/Motivation:
 - *Efficient information discovery in the absence of global knowledge of naming conventions is a fundamental problem in large, decentralized, distributed resource sharing environments such as the Grid*
 - *a document is better described by keywords than by its filename, a computer by a set of attributes such as CPU type, memory, operating system type than by its host name, and a component by its aspects than by its instance name.*
 - *Heterogeneous nature and large volume of data and resources, their dynamism (e.g. CPU load) and the dynamism of the Grid make the information discovery a challenging problem.*
- Key features
 - *P2P system that supports complex queries containing partial keywords, wildcards, and range queries*
 - *Guarantees that all existing data elements that match a query will be found with bounded costs in terms of number of messages and number of nodes involved.*
 - *The system can be used as a complement for current resource discovery mechanisms in Computational Grids (to enhance them with range queries)*



Figure 14

SQUID OPERATION

The overall architecture of SQUID is a distributed hash table (DHT), similar to typical data lookup systems. The key difference is in the way we map data elements to the index space. In existing systems, this is done using consistent hashing to uniformly map data element identifiers to indices. As a result, data elements are randomly distributed across peers without any notion of locality. Our approach attempts to preserve locality while mapping the data elements to the index space. In our system, all data elements are described using a sequence of keywords (common words in the case of P2P storage systems, or values of globally defined attributes - such as memory and CPU frequency - for resource discovery in computational grids). These keywords form a multidimensional keyword space where the keywords are the coordinates and the data elements are points in the space. Two data elements are “local” if their keywords are lexicographically close or they have common keywords. Thus, we map documents that are local in this multi-dimensional index space to indices that are local in the 1-dimensional index space, which are then mapped to the same node or to nodes that are close together in the overlay network. This mapping is derived from a locality-preserving mapping called Space Filling Curves (SFC).

In the current implementation, we use the Hilbert SFC for the mapping, and Chord for the overlay network topology. The overall operation of SQUID is presented in the figure. (a) shows a 2-dimensional keyword space. The data element “*Document*” is described by keywords “*Computer*” and “*Network*”. (b) shows the mapping of the 2-dimensional space to a curve. The query (011, *) defines clusters on the curve (segments). (c) shows the recursive refinement of query (011, *) viewed as a tree. Each node is a cluster, and the bold characters are the cluster's prefixes. (d) illustrates the query resolution process by embedding the leftmost tree path (solid arrows) and the rightmost path (dashed arrows) onto the overlay network topology.

V-GRID AUTONOMIC APPLICATION MANAGEMENT

Truly realistic scientific and engineering simulations require enormous amounts of resources that can surpass even the aggregated capacity of the Grid. The V-Grid (virtual Grid) infrastructure is an application of autonomic computing to science and engineering that is based on the concept of virtualizing grid resources and application execution (analogous to virtual memory). The V-Grid autonomic runtime management framework allows the implementation of a simulation to be driven by the requirements of the science being modeled rather than the size and configuration of the machine that it will be run on.

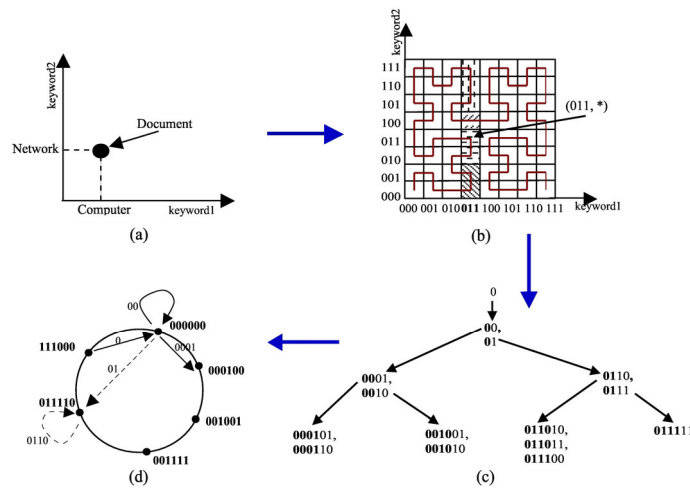
The autonomic behavior in the V-Grid has three primary aspects: (1) V-Grid Monitoring, (2) V-Grid Deduction, and (3) V-Grid Execution.

The V-Grid monitoring engine is a decentralized entity composed of context agents that provides application and system context awareness. Application monitoring uses sensors exported by the autonomic components and services and provides information about the current state, dynamics and requirements of components and the application. System/resource monitoring builds on context information provided by OGSA and existing Grid middleware (e.g. NWS, Globus, Autopilot) and extends their capabilities to support dynamic monitoring requirements and information aggregation.

The V-Grid deduction engine uses application/components specifications, context and predicted behavior to deduce objective functions and execution and management strategies. This includes identifying and characterizing natural regions, defining Virtual Computational Units or VCU that reflect the current state of the application, mapping them onto Virtual Resource Units or VRUs based on their specifications, and outlining scheduling policies and constraints. This mapping of VCUs onto VRUs exploits the spatial, temporal and functional heterogeneity of the application to reduce couplings and maximize performance.

The V-Grid execution engine implements policies and strategies defined by the deduction engine using OGSA and autonomic Grid services. The main activities of this engine are (1) dynamic reservation and allocation of VRUs, (2) adaptive mapping and scheduling of VCUs to VRUs, and VRUs to physical resources, and (3) autonomic management, control and adaptation of application execution.

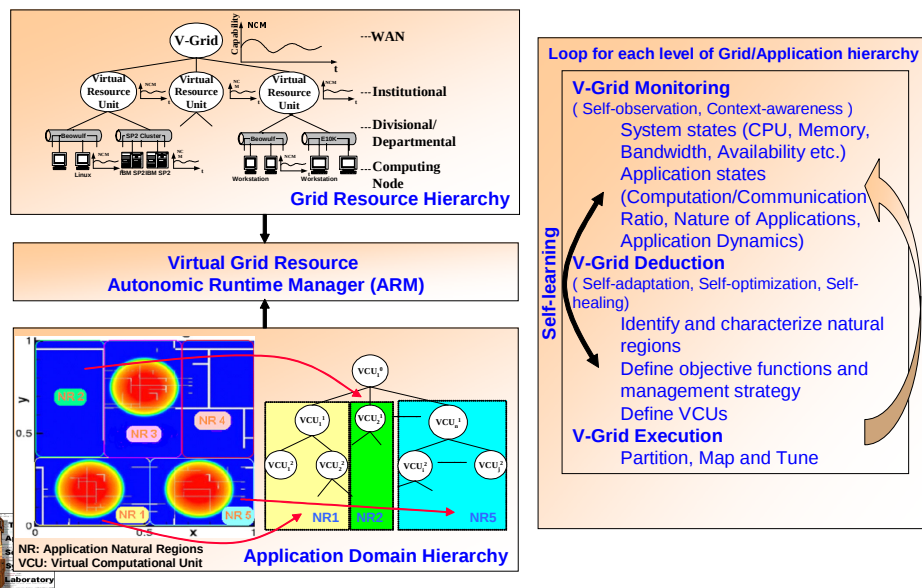
SQUID: Operation



T
A
S
S
L
The
Applied
Software
Systems
Laboratory

Figure 15

V-Grid: Autonomic Application Management



T
A
S
S
L
The
Applied
Software
Systems
Laboratory

Figure 16

ADAPTIVE MESH REFINEMENT

Dynamically adaptive mesh refinement (AMR) methods for the numerical solution to partial differential equations (PDEs) employ locally optimal approximations, and can yield highly advantageous ratios for cost/accuracy when compared to methods based upon static uniform approximations. These techniques seek to improve the accuracy of the solution by dynamically refining the computational grid in regions with large local solution error.

Structured AMR (SAMR) techniques start with a coarse base grid with minimum acceptable resolution that covers the entire computational domain. As the solution progresses, regions in the domain with large solution error, requiring additional resolution, are identified and refined. Refinement proceeds recursively so that the refined regions requiring higher resolution are similarly tagged and even finer grids are overlaid on these regions. The resulting grid structure is a dynamic adaptive grid hierarchy (such as the SAMR formulation by Berger and Oliger, shown in the figure).

Methods based on SAMR can lead to computationally efficient implementations as they require uniform operations on regular arrays and exhibit structured communication patterns. Distributed implementations of these methods, however, lead to interesting challenges in dynamic resource allocation, data-distribution, load-balancing, and runtime management.

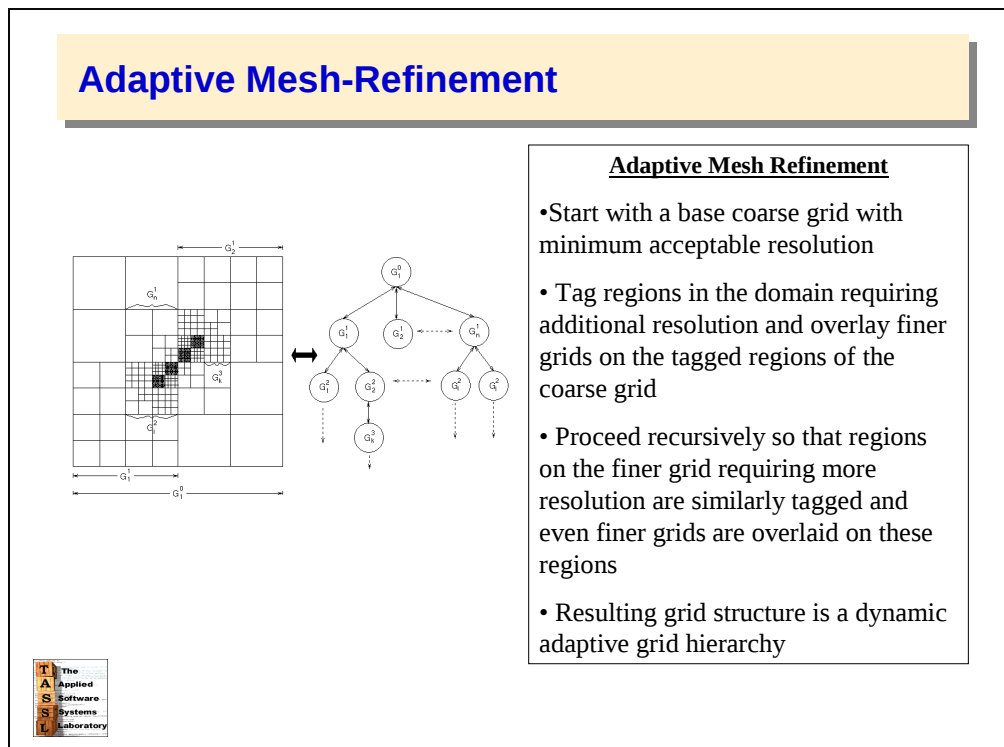


Figure 17

STRUCTURE ADAPTIVE MESH REFINEMENT APPLICATIONS

Structured adaptive mesh refinement (SAMR) methods are being effectively used for adaptive PDE solutions in many domains, including computational fluid dynamics, numerical relativity, astrophysics, and subsurface modeling and oil reservoir simulation.

The top-left application belongs to the Zeus kernel coupled with GrACE (SAMR infrastructure) and Cactus (problem solving environment) packages, and shows a 3-D blast wave in the presence of a uniform magnetic field with 3 levels of refinement. Zeus-MP solves the equations of ideal (non-resistive), non-relativistic, hydrodynamics and magnetohydrodynamics, including externally applied gravitational fields and self-gravity.

The top-right figure is taken from the IPARS oil reservoir simulator and shows the multi-block grid structure and oil concentration contours. The MACE (Multi-block Adaptive Computational Engine) infrastructure support multi-block grids where multiple distributed and adaptive grid blocks with heterogeneous discretization are coupled together with lower dimensional mortar grids.

The CCA (Common Component Architecture) and GrACE application at bottom-left investigates the direct numerical simulation of flames with detailed chemistry solving the Navier-Stokes and species evolution equations without approximations. The figure shows this simulation for a mixture of H_2 and Air in stoichiometric proportions, with 3 hot spots at 1000K causing H_2 -Air mixture to ignite and create many different radicals. The scientific problems being studied are the flame stabilization mechanisms of unsteady laminar and turbulent flames, with emphasis on the flame structure at the flame base.

The bottom-right application simulates the dynamic response of materials, with the goal to develop a Virtual shock physics Test Facility (VTF) for a wide range of compressive, tensional, and shear loadings, including those produced by detonation of energetic materials. GrACE is the computational engine underlying the VTF. The figure shows the compressible turbulence simulation solving the Richtmyer-Meshkov instability in 3D (RM3D) using adaptive refinements. The Richtmyer-Meshkov instability is a fingering instability that occurs at a material interface accelerated by a shock wave.

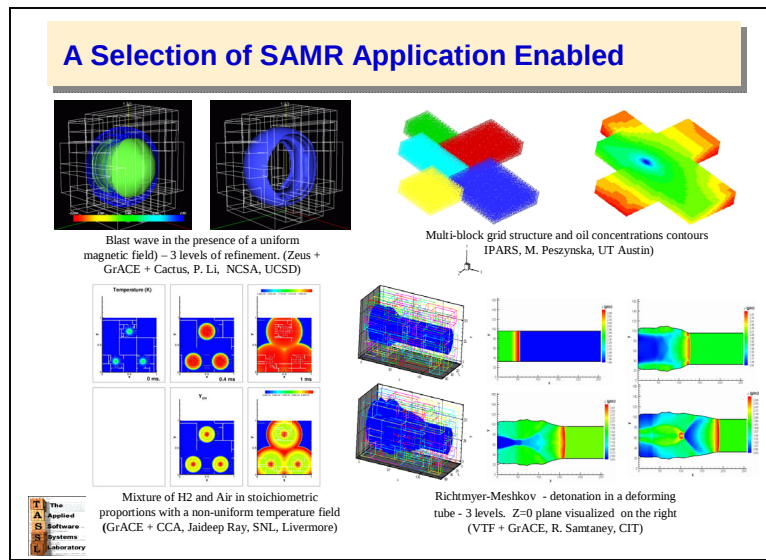


Figure 18

ARMADA: AUTONOMIC RUNTIME MANAGEMENT OF DYNAMIC APPLICATIONS

ARMaDA is a framework for the autonomic run-time management and optimization for dynamic SAMR applications. Autonomic behavior is achieved by adapting SAMR application execution to optimize partitioning, load-balancing, and scheduling. Adaptation parameters include the partitioning scheme based on current runtime state (GrACE, Vampire, etc.), granularity/patch size affecting load balance and overhead, dynamic allocation of processors (from beginning or “on-demand”). Other optimizations include hierarchical decomposition using dynamic processor groups, communication optimization, latency tolerance, multithreading, etc.

Autonomic application management involves system-sensitive and application-sensitive adaptation. System-sensitive application management uses current and predicted system state characterization to make application adaptation decisions. For example, the information about the current load and available memory may determine the granularity of the mapping of the application components to the processing nodes, while the availability and “health” of the computing elements on the grid may determine the nature (refined grid size, aspect ratios, etc.) of refinements to be allowed.

Application sensitive adaptations use the current state of the application to drive the run-time adaptations. The abstraction and characterization of the application state is used to drive the resource allocations, partitioning and mapping of application components onto the grid, selection of partitioning and load-balancing algorithms and their configurations, communication mechanisms, etc.

ARMaDA: Autonomic Run-time Management and Optimization for Dynamic (SAMR) Applications

- Partitioning, load-balancing and scheduling of SAMR applications.
 - *Partitioning Scheme*
 - “Best” partitioning based on application/system configuration and current application/system state
 - G-MISP+SP, pBD-ISP, SFC (Vampire, GrACE, Zoltan, ParMetis, ...)
 - *Granularity*
 - patch size, AMR efficiency, comm./comp. ratio, overhead, node-performance, load-balance, ...
 - *Number of processors/Load per processor*
 - Dynamic allocations/configuration/management
 - 1000+ processor from the beginning or “on-demand”
 - *Hierarchical decomposition using dynamics processor groups*
 - *Communication optimizations/latency tolerance/multithreading*
 - *Availability, capabilities, and state of system resources*
 - SNMP, NWS



Figure 19

ARMADA: AUTONOMIC RUNTIME MANAGEMENT

Starting in the upper-left of the figure, the SAMR application is monitored by the V-Grid Monitoring Engine to enable the V-Grid Planning and Analysis Engines to identify natural regions and characterize application state. Simultaneously, the V-Grid Monitoring Engine also monitors and characterizes the system. The synthesized system capability combines monitored information with history and predictive models. Both of these characterizations flow into the V-Grid Analysis and Execution Engines. The V-Grid Analysis Engine deduces objective functions, strategies, and normalized work and resource metrics, using policies and constraints to navigate the decision space. The V-Grid Execution engine uses this information to autonomously partition or repartition the application into VCUs that are mapped and scheduled onto VRUs. Global-Grid Scheduling (GGS) is first used across VRUs and then Local-Grid Scheduling (LGS) within a VRU. The V-Grid Execution Engine then allocates and configures Grid resources and schedules execution of VRUs. This execution is, in turn, is monitored by the monitoring engine. This flow of events occurs within a distributed framework.

A dynamic topology of V-Grid framework agents will locally monitor the application and resources. Changes in the local natural regions will be monitored along with changes in the local resource performance. The V-Grid Analysis Engine may be able to make many local decisions, but may also be able to make improved decisions by “comparing notes” with neighboring framework agents. The autonomic partitioning and scheduling may move work among agents or may acquire new resources and add new agents to the framework.

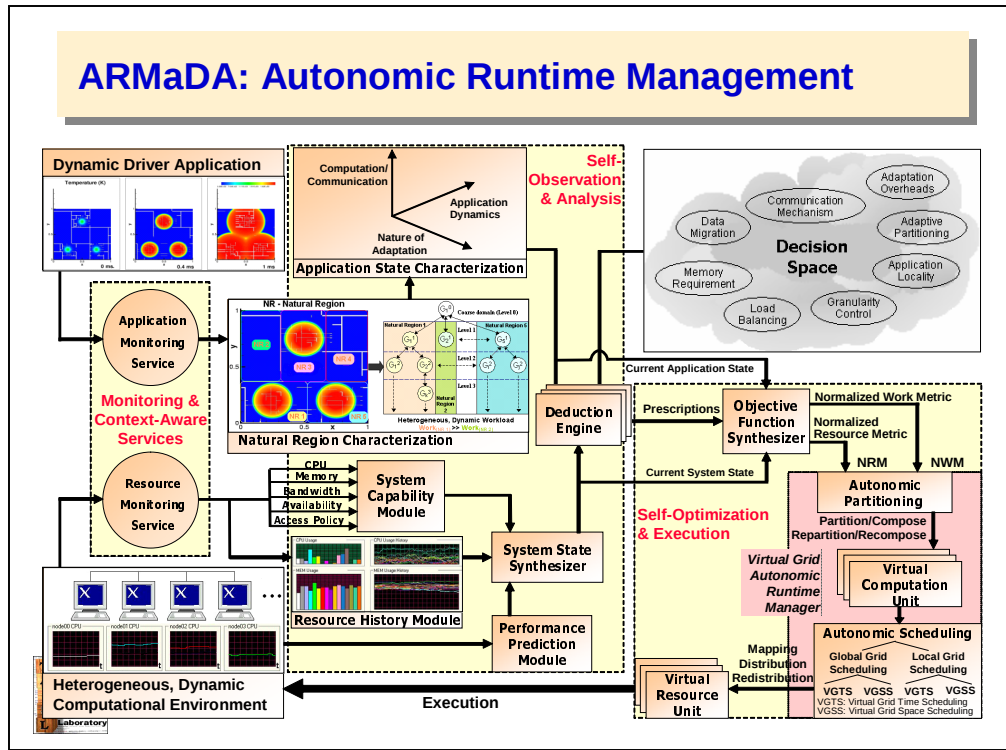


Figure 20

ARMADA: APPLICATION-SENSITIVE ADAPTATIONS

The ARMaDA framework performs adaptive application-sensitive partitioning based on the input parameters and the application's current runtime state. Partitioning behavior is characterized based on the {Partitioner, Application, Computer System} (PAC) tuple. Each PAC tuple is evaluated using a 5-component metric that includes load imbalance, communication requirement, amount of data migration, partitioning induced overhead, and the partitioning time. The PAC relationship is dynamic and the partitioner P is a function of the state of the application A and the computer system C at that time. The octant approach is used to classify application runtime state with respect to the adaptation pattern, computations/communications, and activity dynamics.

The ARMaDA framework has three components: application state monitoring and characterization, partitioner repository and policy engine, and an adaptation component. The state characterization component implements mechanisms that abstract the current application state in terms of the computation/communication requirements, application dynamics, and the nature of the adaptation. The policy engine provides an association for mapping octants to partitioners and the partitioning repository includes a selection from popular software tools such as GrACE (ISP) and Vampire (pBD-ISP, GMISP+SP). Subsequently, the meta-partitioner or adaptation component dynamically selects the appropriate partitioner at runtime and configures it with associated parameters such as granularity. As shown in the slide, experimental results demonstrate the improvement in SAMR application execution using application-sensitive partitioning – 26.19% for VectorWave-2D application on 32 processors on Linux Beowulf cluster “Frea” and 38.28% for RM3D application on 64 processors on IBM SP2 “Blue Horizon”.

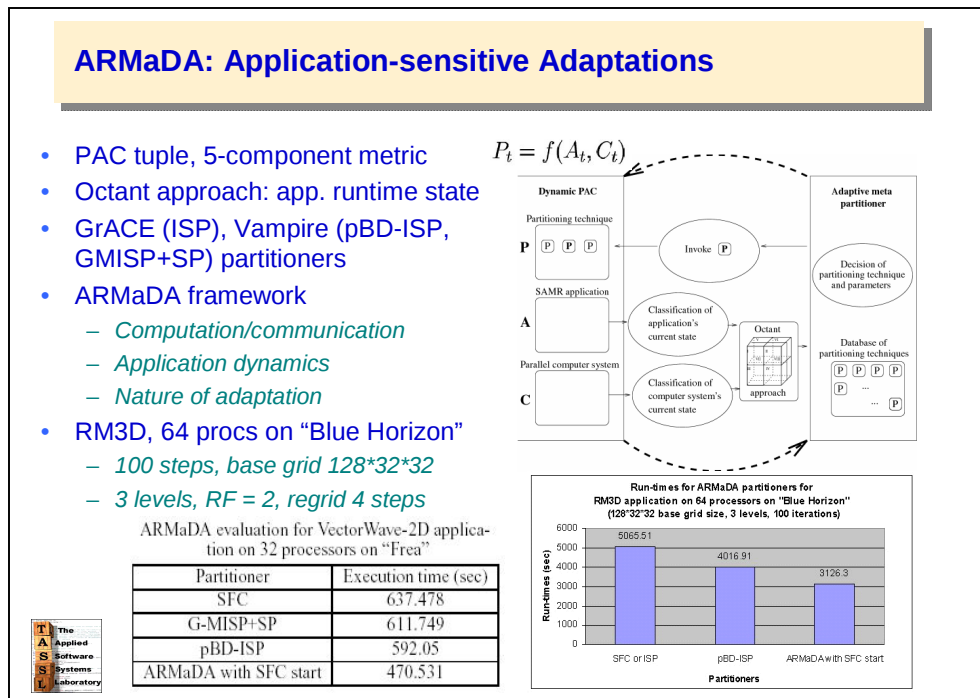


Figure 21

ARMADA: SYSTEM-SENSITIVE ADAPTATIONS

The ARMaDA framework reacts to system capabilities and current system state to select and tune distribution parameters by dynamically partitioning and load balancing the SAMR application grid hierarchy. Current system state is obtained at runtime using the Network Weather Service (NWS) resource monitoring tool. NWS measurements include CPU availability, end-to-end network bandwidth, free memory, and the amount of space unused on a disk. System state information along with system capabilities are then used to compute the relative capacity of each computational node as a weighted sum of the normalized system metric. The weights are application dependent and reflect its computational, memory, and communication requirements. These relative capacities are used by the “system-sensitive” partitioner for dynamic distribution and load-balancing.

The system-sensitive partitioner is evaluated using the RM3D CFD kernel on a 32-node Linux-based workstation cluster. The kernel used 3 levels of factor 2 space-time refinements on a base mesh of size 128*32*32. System-sensitive partitioning reduced execution time by about 18% in the case of 32 nodes. The table in the slide illustrates the effect of sensing frequency on overall application performance. Dynamic runtime sensing improves application performance by as much as 45% compared to sensing only once at the beginning of the simulation. In this experimental setup, the best application performance was achieved for a sensing frequency of 20 iterations.

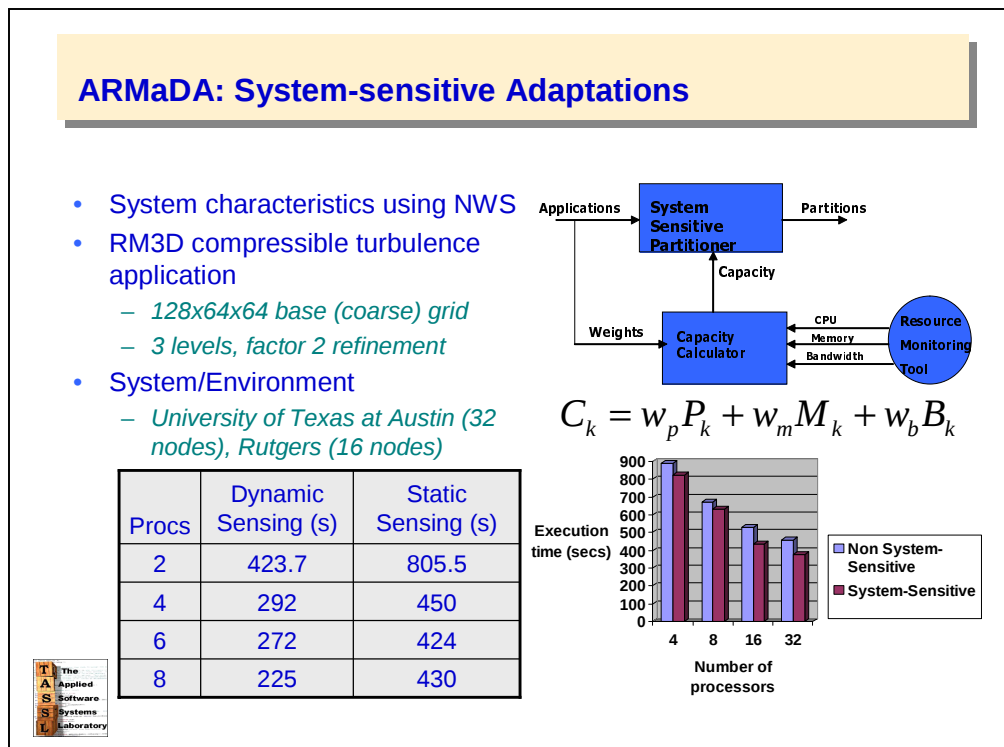


Figure 22

ARMADA: PROACTIVE MANAGEMENT

The ARMaDA framework uses performance prediction functions to estimate execution time and application performance. Performance Functions (PF) describe the behavior of a system component, subsystem or compound system in terms of changes in one or more of its attributes. The PFs of each resource used by an application can be composed to generate an overall end-to-end PF that quantifies application performance.

Performance functions model the application execution time for SAMR-based RM3D and describe overall behavior with respect to the computational load metric on the machine of choice (such as IBM SP “Seaborg” and Linux Beowulf “Discover”). The evaluation on IBM SP yields 2 PFs for small loads ($\leq 30,000$ work units) and large loads ($> 30,000$ units) respectively, whereas the Linux Beowulf produces a single PF. The error in modeling the execution time is low – 0-8% for IBM SP and 0-6% for Linux Beowulf.

The PF modeling approach is used by the ARMaDA framework to determine when the benefits of dynamic load redistribution exceed the costs of repartitioning and data movement (if workload imbalance exceeds a certain threshold). A threshold of 0 indicates regular periodic load redistribution while a high threshold represents the ability of the application hierarchy to tolerate workload imbalance. The RM3D evaluation on 8 processors on Linux Beowulf cluster analyzes the effect of dynamic load redistribution on application recompose time for redistribution thresholds of 0 and 1. The application uses 3 refinement levels on a base mesh of size $64 \times 16 \times 16$ with regridding every 4 steps. Threshold of 1 considers the costs of redistributing load and results in recompose time being reduced by half (improvement of almost 100%) as compared to a threshold of 0.

ARMaDA: Proactive Management

- Performance Function (PF) – behavior in terms of attribute changes
- “Computational load” metric to model RM3D execution time
- IBM SP “Seaborg” (NERSC)
 - PF_s – small loads (≤ 30000 units), PF_h – large loads (> 30000)
 - Error in modeling execution time is low (0 – 8%)
- Linux Beowulf “Discover” (Rutgers)

- Single PF
 - Error in modeling execution time is low (0 – 6%)
$$PF = \sum_{i=0}^{10} b_i * x^i$$
- Dynamic load redistribution for RM3D & effect on “recompose” time
 - 8 processors, base mesh $64 \times 16 \times 16$, 3 levels of factor 2 refinements
 - Redistribution thresholds of 0 and 1
 - Thresh=1 improves recompose time by 100% compared to thresh=0



Figure 23

AUTONOMIC OIL WELL PLACEMENT

The goal of this application is to dynamically optimize the placement and configuration of oil wells to maximize revenue. The peer components involved include:

1. Integrated Parallel Accurate Reservoir Simulator (IPARS) providing sophisticated simulation components that encapsulate complex mathematical models of the physical interaction in the subsurface, and execute on distributed computing systems on the Grid.
2. IPARS Factory responsible for configuring IPARS simulations, executing them on resources on the Grid and managing their execution.
3. Very Fast Simulated Annealing (VFSA) optimization service based on statistical physics and the analogy between the model parameters of an optimization problem and particles in an idealized physical system.
4. Economic Modeling Service that uses IPARS simulation outputs and current market parameters (oil prices, costs, etc.) to compute estimated revenues for a particular reservoir configuration.
5. Discover Middleware that integrates Globus Grid services (GSI, MDS, GRAM, and GASS), via the CORBACog, and Discover remote monitoring, interactive steering, and collaboration services, and enables resource discovery, resource allocation, job scheduling, job interaction and user collaboration on the Grid.
6. Discover Collaborative Portals providing experts (scientists, engineers) with collaborative access to other peer components. Using these portals, experts can discover and allocate resources, configure and launch peers, and monitor, interact with, and steer peer execution. The portals provide a shared workspace and encapsulate collaboration tools such as Chat and Whiteboard.

(This slide is courtesy M. Peszynska)

Autonomic Oil Well Placement

- Optimization algorithm: use VFSA (Very Fast Simulated Annealing)
 - *requires function evaluation only, no gradients*
- IPARS delivers
 - *fast-forward model (guess->objective function value)*
 - *post-processing*
- Formulate a parameter space
 - *well position and pressure (y,z,P)*
- Formulate an objective function:
 - *maximize economic value $Eval(y,z,P)(T)$*
- Normalize the objective function $NEval(y,z,P)$ so that:

$$\min Neval(y,z,P) \Leftrightarrow \max Eval(y,z,P)$$



Figure 24

AUTONOMIC OPTIMIZATION OF OIL RESERVOIR

These peer entities involved in the optimization process need to dynamically discover and interact with one another as peers to achieve the overall application objectives. The experts use the portals to interact with the Discover middleware and the Globus Grid services to discover and allocate appropriate resource, and to deploy the IPARS Factory, VFSA and Economic model peers ((1)). The IPARS Factory discovers and interacts with the VFSA service peer to configure and initialize it ((2)). The expert interacts with the IPARS Factory and VFSA to define application configuration parameters ((3)). The IPARS Factory then interacts with the Discover middleware to discover and allocate resources and to configure and execute IPARS simulations ((4)).

The IPARS simulation now interacts with the Economic model to determine current revenues, and discovers and interacts with the VFSA service when it needs optimization ((5)). VFSA provides IPARS Factory with optimized well information ((6)), which then launches new IPARS simulations ((7)). Experts at anytime can discover and collaboratively monitor and interactively steer IPARS simulations, configure the other services and drive the scientific discovery process ((8)). Once the optimal well parameters are determined, the IPARS Factory configures and deploys a production IPARS run.

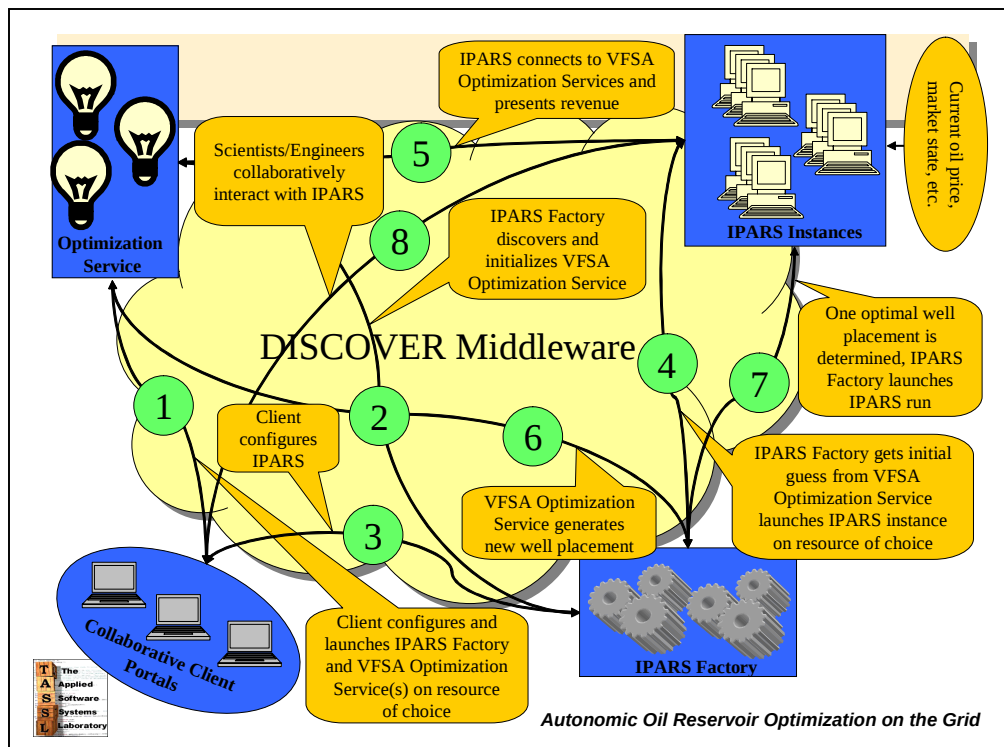


Figure 25

AUTONOMIC OIL WELL PLACEMENT

The figure below show results from the autonomic oil well placement applications. It shows that the process converges to the optimal placement in 20 iterations.

(This slide is courtesy M. Peszynska)

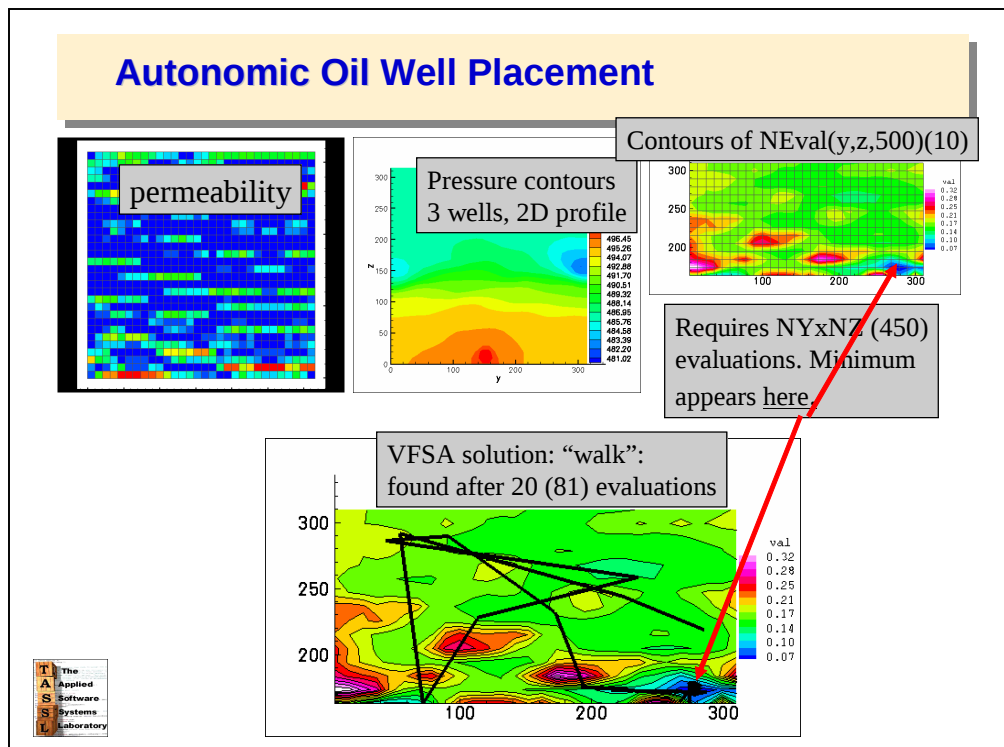


Figure 26

CONCLUSION

The computational solutions addressed by the AutoMate project are based on fundamental innovations in the development, optimization and deployment of component-based Grid applications, thereby allowing the heterogeneity and dynamics of the applications to match that of the Grid and fully exploit its potential. These innovations will enable scientists to choreograph high performance, integrated end-to-end simulations that were never possible or attempted before. The key IT contributions are the methodology and associated technologies that enable the development of applications that can manage and exploit the dynamism and heterogeneity of the Grid, and that address the extremely serious problem of software complexity that is threatening both academia and industry.

We currently have working prototypes of each of the components presented in this paper, and are in the process of integrating them to support autonomic structured adaptive mesh refinement applications (SAMR) in science and engineering. Further information about AutoMate and its components can be obtained from <http://automate.rutgers.edu>.

Conclusion

- Autonomic (adaptive, interactive) applications can enable accurate solutions of physically realistic models of complex phenomenon.
 - *their implementation and management in Grid environments is a significant challenge*
- AutoMate provides key technologies to enable the development of autonomic Grid applications
 - *ACCORD: Autonomic application framework*
 - *RUDDER: Decentralized deductive engine*
 - *SESAME: Dynamic access control engine*
 - *Pawn: P2P messaging substrate*
 - *SQUID: P2P discovery service*
- Application scenarios
 - *V-Grid autonomic runtime management of SAMR applications*
 - *Autonomic optimization of oil reservoirs*
- More Information, publications, software
 - www.caip.rutgers.edu/TASSL/Projects/AutoMate/
 - automate@caip.rutgers.edu / parashar@caip.rutgers.edu



Figure 27

**THE NEXT WAVE OF UBIQUITOUS COMPUTING IN
KNOWLEDGE ECONOMY: CHALLENGES AND
OPPORTUNITIES**

Youngjin Yoo
Case Western Reserve University
Cleveland, OH

THE NEXT WAVE OF UBIQUITOUS COMPUTING IN KNOWLEDGE ECONOMY: CHALLENGES AND OPPORTUNITIES

In this presentation, I will focus on the organizational opportunities and challenges that ubiquitous computing brings to organizations. Whether organizations like it or not, a fundamental paradigm shift in organizational computing is taking place. This, along with changes in the society and economy in general, presents new opportunities and challenges to organizations that they've never faced before.

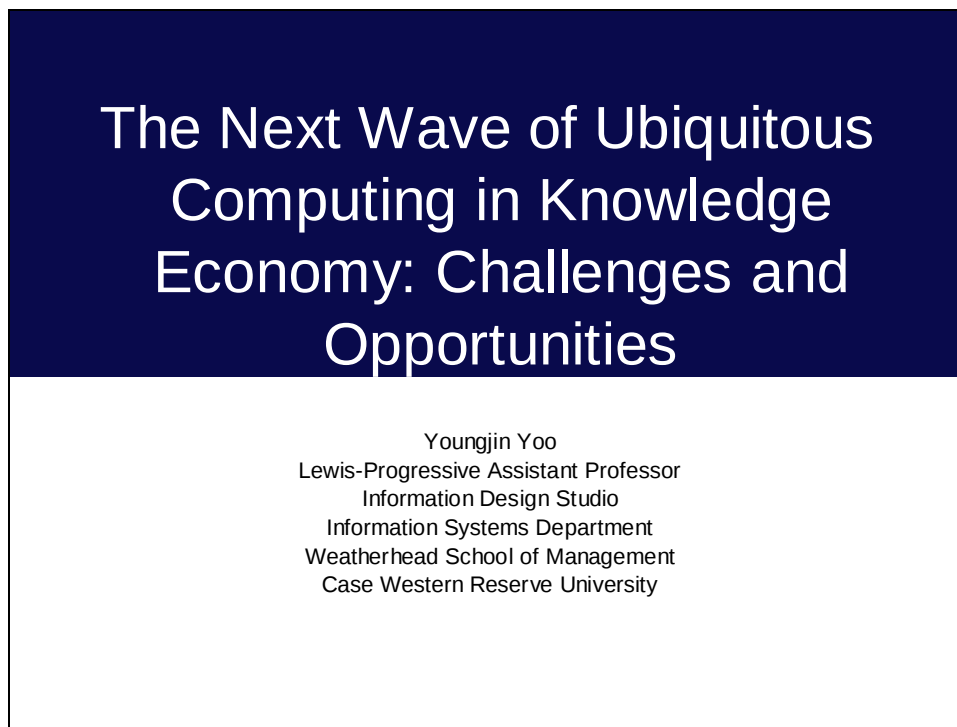


Figure 1

FOUR WAVES OF ORGANIZATIONAL COMPUTING

We can think of four distinctive waves of computing paradigms in organizations. It starts with data processing era where mainframe computers were used to automate the back-office tasks. In late 70's and early 80's, this wave was replaced by Micro wave, represented by personal computers and end-user productivity software (Word Perfect, Lotus 1-2-3, dBase III+, and Harvard Graphics). The introduction of local area network and, later, the Internet, once again changed the nature of computing and took us into Network era. At this point, we are yet again experiencing a transition from network to ubiquitous wave. Each wave of computing not only represents more advanced and powerful computer hardware and software, but also changes of the strategic significance of IT in organizations. It is journey from the back-office to front-office. It is a journey from being utility to strategic assets. This trend will continue in ubiquitous wave.

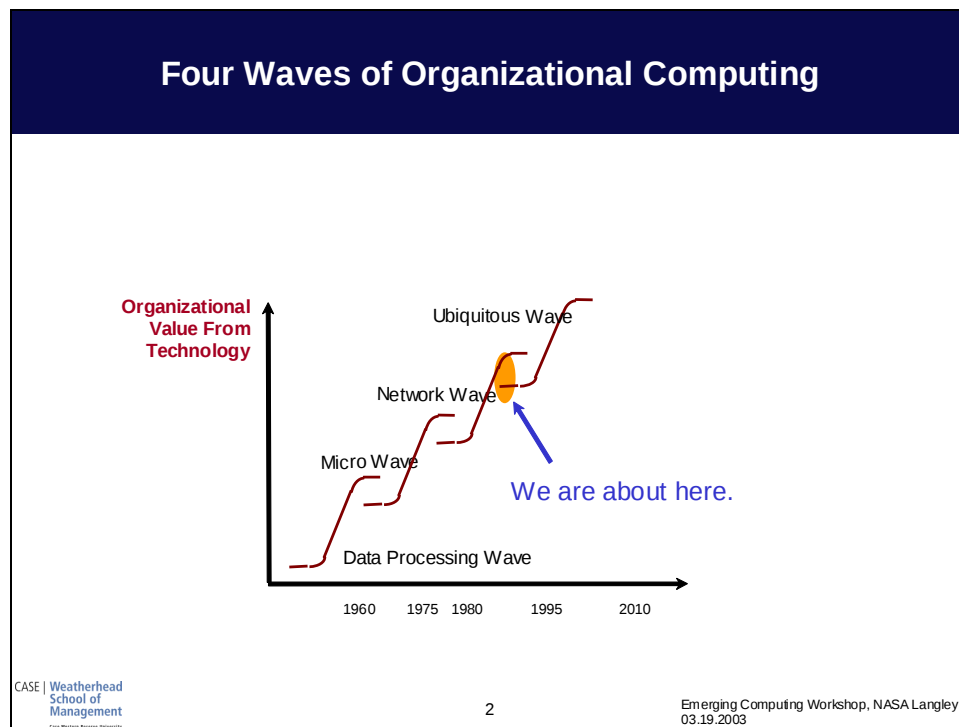


Figure 2

CHANGES IN BUSINESS ENVIRONMENTS

In order to properly understand the importance of ubiquitous computing, we need to put it into a context of the current societal and economical environments. First, it is knowledge economy where knowledge is the primary means to add and create values. In knowledge economy, being connected is more important than having possessions. In the industrial economy, physical products were to be purchased and owned in order to consume. In knowledge economy, consumers need to be connected to experience knowledge-based services. As such, physical assets are not as important as knowledge assets. Second, it is networked economy. The value is created not by a heroic individual or a single firm, but rather created by a community of distributed agents. This requires a fundamental shift in our thinking about organizing. Finally, we are facing a fundamentally different market with customers who grew up with computers and Nintendo video games. These global new generations of customers emphasize aesthetics and spontaneity in their consumer experiences.

Changes in Business Environments

- Knowledge economy
 - ✓ From possession to connection
 - ✓ From having to experiencing
 - ✓ From physical products to knowledge products
 - ✓ From physical assets to knowledge resources
 - ✓ Values are created through the integration of knowledge resources
- Networked economy
 - ✓ From a single firm to a pack of runners
 - ✓ From an individual hero to a community of distributed agents
 - ✓ From stand-alone machines to socio-technical web
- Different market
 - ✓ Nintendo generation
 - ✓ Spontaneity
 - ✓ Global

CASE | Weatherhead
School of
Management

3

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 3

AN ENVIROMENT IN CHANGE

Such technological and environmental changes, along with series of recent de-regulations, have created a new environment. In this new environment, traditional separate industries come together and compete in the same space.

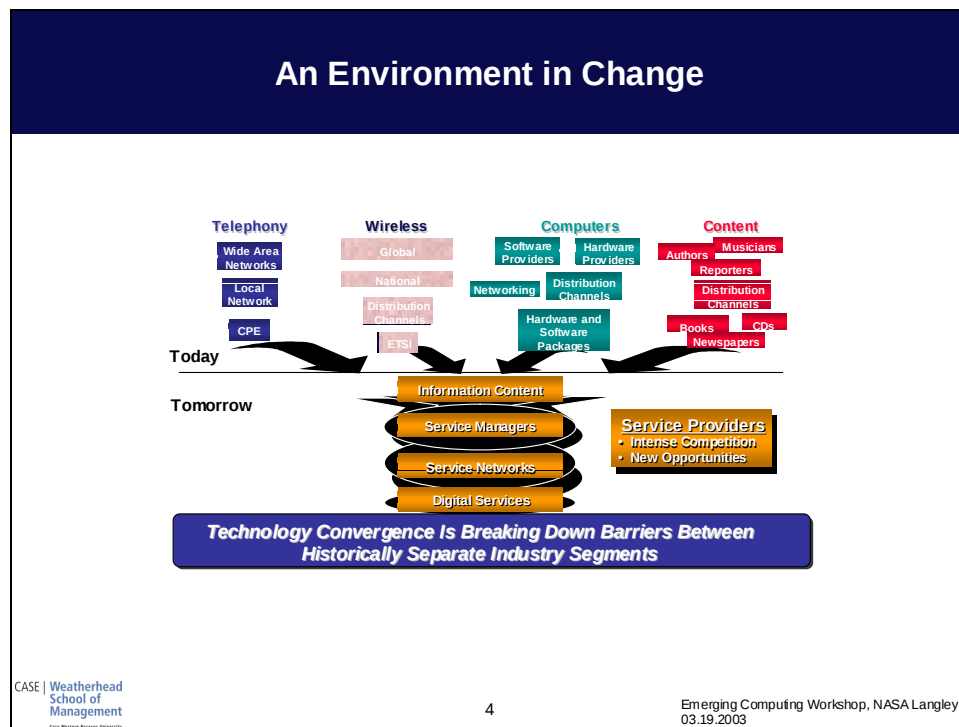


Figure 4

DRIVERS FOR DIGITAL SERVICES

These changes in the economy, society, and technology can be summarized as the emergence of digital services as core elements of economic activities in the economy. Such emergence of digital service is not just technical nor organizational. It is socio-technical shift in the society. This is far more significant and fundamental than the emergence of web-commerce. In fact, much of the prior technological innovations in organizations (such as e-commerce, business process reengineering, enterprise resource planning systems, etc) can be seen as fundamental basis for this unavoidable emergence of digital service economy.

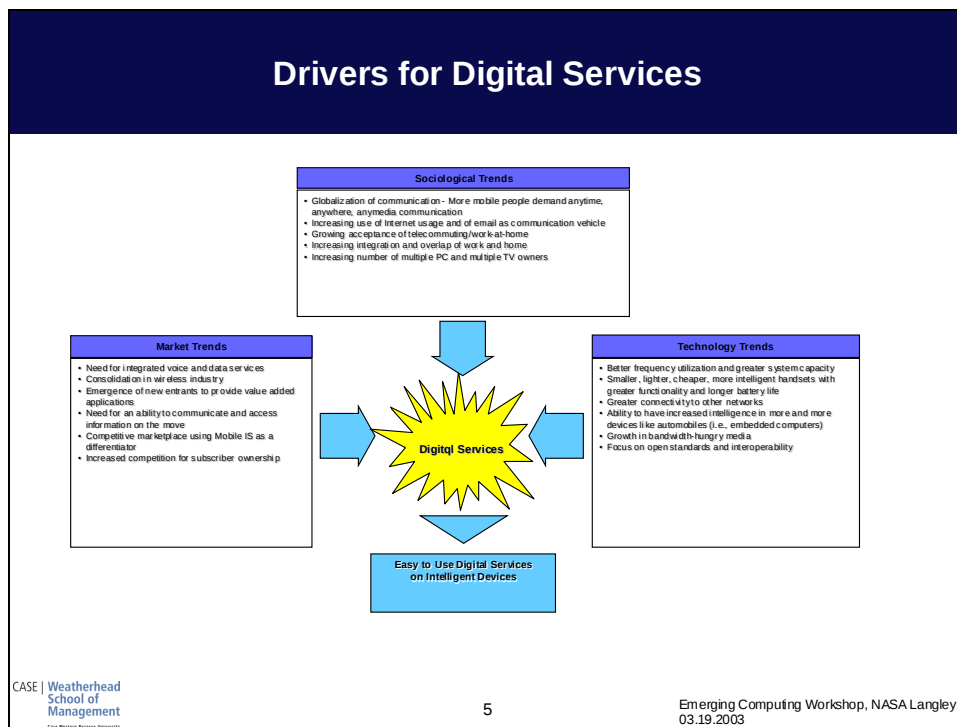


Figure 5

THE NEW DIGITAL ENVIRONMENT

Thus, in this new digital environment, there will be both technological push and market pull. While the emergence of knowledge economy demands the anytime, anyplace delivery service, the technology will be there to enable such ubiquitous digital business transactions. Similarly, digital convergence will enable mass customization.

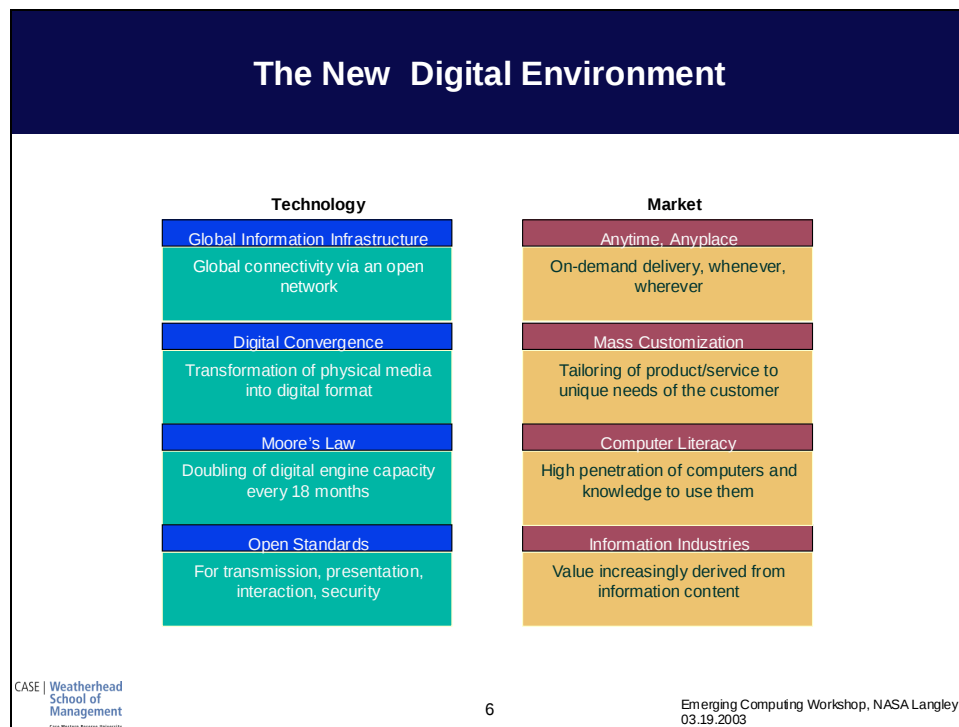


Figure 6

TECHNOLOGICAL CHANGES

All of these changes in the economy and society are fueled further by the relentless developments in technology in all areas. In particular, the next wave of organizational computings will be built on heterogeneous distributed computing infrastructure along with novel technology kernels in hardware, software standards, and telecommunication and network technologies. In the next few slides, I will examine three key drivers of this new technological environment.

Technological Changes

- Fast change in all computing technologies
- Heterogeneous and distributed computing
 - ✓ Novel technology kernels (hardware, telecom, system software- standards)
 - ✓ Distributed system architectures (design, control, performance)
 - ✓ Heterogeneous interoperability (services, semantics, metadata, ontologies)
 - ✓ Key features: mobility, net-centric services, intelligent agents



CASE | Weatherhead
School of
Management
www.case.edu

7

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 7

THREE KEY TRENDS

In short, the new organizational computing environments can be characterized with three key words: mobility, digital convergence, and mass scale.

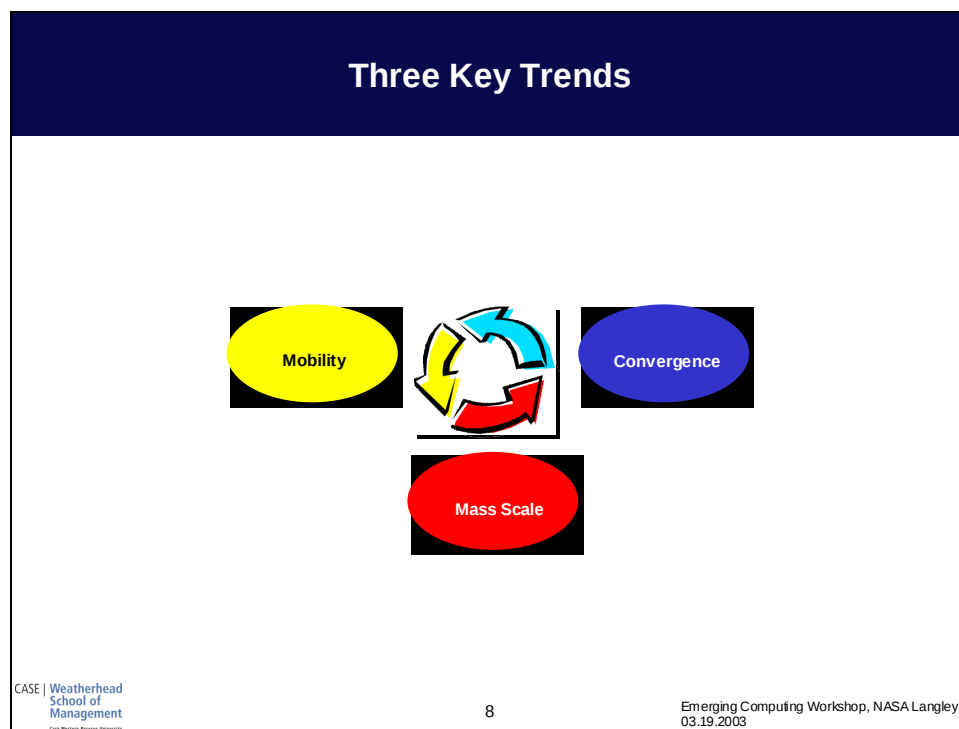


Figure 8

DIGITAL CONVERGENCE

The rapid penetration of digital technology in all form of computing and communication have been enabled by the dramatic reduction of the computing costs and the emergence of open standards and new chip designs. The digital convergence enables new forms of engagement with digital services and new services such as in entertainment and telematics areas. Often, these new services require integrations of services traditionally offered in separate channels.

Digital convergence

- **Digital convergence:** enabled by computing costs and chip design + open standards
 - ✓ New forms of engagement with digital services
 - ✓ New services (entertainment, telematics)
 - ✓ Integration of services (video+ data) Challenges
- **Challenges**
 - ✓ Requires independence between the content and the medium (Ex: CNN service)
 - ✓ Requires miniaturization of devices

CASE | Weatherhead School of Management

9

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 9

MOBILITY

The second major characteristics is the mobility. We often focus on physical mobility as uni-dimensional concept. However, mobility can be divided into micro-, local-, and remote-mobility depending on the geographical coverage on the area that is covered. We also need to think of social mobility as well as physical mobility. In the past where particular computing activities were tied to particular time-space combination, social mobility was relatively stable--one social role in one physical place. However, as physical mobility of computing devices become higher, one can have a high degree of social mobility even within the same geographic location and temporal boundary. In order to support both physical and social mobility, organizations need to develop socio-technical ontology.

Mobility

- **Mobility** covers physical mobility and social mobility
 - ✓ Social: roles, capabilities, rights, preferences
 - ✓ Physical: micro mobility, local mobility, remote mobility
 - ✓ Requires **mobility of services** across platforms
- Enables new services as combinations of social and physical mobility and independence between services and locations
- Challenges
 - ✓ **Interoperability** and **peer-to-peer synchronization** becomes critical
 - ✓ Requires ☐ ☐ ☐ **social ontology** to support social mobility

CASE | Weatherhead School of Management

10

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 10

MASS SCALE

The combination of digital convergence and mobility lead to unprecedented large scale of deployment of various types of digital services at a global level. In this new mass scale era, new technical challenges emerge including scalability, reliability, complexity, security, and performance. New emergence of grid computing and autonomic computing, for example, will play vital roles to support such a mass scale.

Mass Scale

- **Mass scale:** services provided in principle at a global level, pervasiveness implies high volumes
 - ✓ Internet capable mobile devices: 1 billion by 2003
 - ✓ 300 million Bluetooth devices in US alone by 2003
 - ✓ PDA sales in US in 2000 was \$1.03 billion
- Challenges: scalability, reliability, complexity, security and performance
- These are affected by both mobility (coverage, network features) and digital convergence (bandwidth, QoS)

CASE | Weatherhead
School of
Management

11

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 11

UBIQUITOUS INFORMATION ENVIRONMENT

To summarize, a ubiquitous computing environment can be understood as a heterogeneous assemblage of interconnected technological and organizational elements, enabling both physical and social mobility of computing communication services between organizational actors both within and across organizational boundaries. The impact and challenges of ubiquitous computing need to be understood as an integral part of modern complex organizations as socio-technical webs of distributed intelligent agents.

Ubiquitous Information Environment

- A heterogeneous assemblage of interconnected technological and organizational elements, enabling physical and social mobility of computing and communication services between organizational actors both within and across organizational boundaries
- The Theme of *International Conference on Information Systems 2003*: IT Everywhere
- An integral part of modern complex organizations as socio-technical webs of distributed intelligent agents

CASE | Weatherhead School of Management

12

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 12

MOVEMENTS TO UBIQUITOUS COMPUTING

This figure shows the movement from traditional organizational computing where both mobility and the degree of embeddedness of computing in environments were low to ubiquitous computing where both of them are high. Further, it shows the conceptual differences among pervasive computing, mobile computing, and ubiquitous computing.

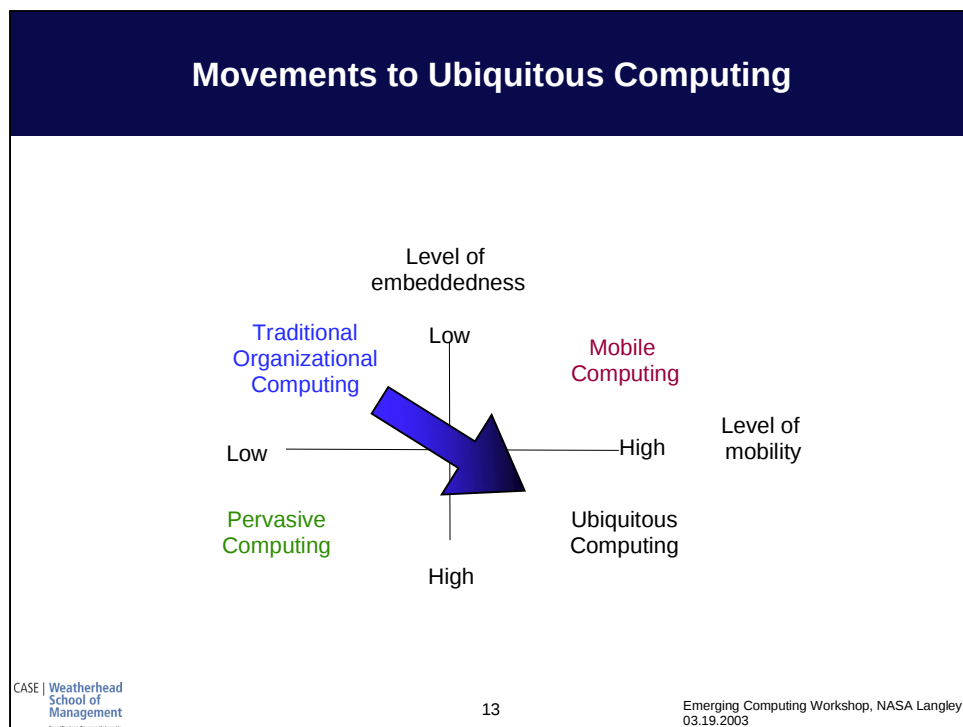


Figure 13

OPPORTUNITIES

These technological developments can enable new and disruptive business models and opportunities. For example, Progressive insurance experimented with “usage-based” auto insurance policy premium model in Texas. The service was enabled by GPS chips with a dial-up modem installed in cars along with powerful database of past history. Customers were charged based on their actual driving patterns, rather than their personal profiles. Such a revolutionary product enabled by the combination of powerful ubiquitous computing tools can potentially cause dramatic disruption in the market. Organizations need to proactively seek to leverage this emerging ubiquitous computing tools in order to create this type of disruptive opportunities.

Opportunities

- New and disruptive business models are possible
- Examples
 - ✓ Telematics
 - ✓ Home digital media
 - ✓ High-velocity coordination systems
 - ✓ On-demand distributed training and learning

CASE | Weatherhead
School of
Management

14

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 14

CHALLENGES

However, organizations need to overcome significant technical, organizational institutional challenges in order to take advantage of emerging ubiquitous computing.

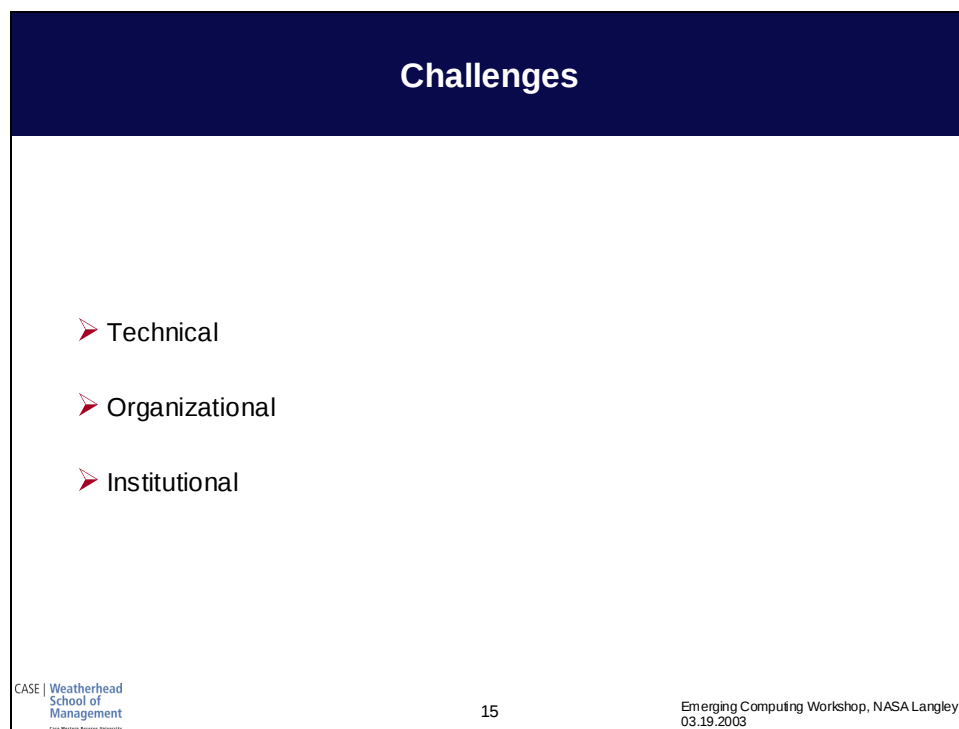


Figure 15

INFRASTRUCTURE

Technical challenges can be divided into two: infrastructure and services. For infrastructure, ubiquitous computing means technically heterogeneous, geographically dispersed, and institutionally complex without centralized coordination mechanism. Thus, as pointed out earlier, providing interoperability, scalability, stability, reliability and persistence through infrastructure will be critically important. Furthermore, since ubiquitous computing involves many diverse devices, seamless integrations among different devices, services and platforms will become key IT management issues. The ubiquitous computing infrastructure need to provide location awareness, service availability, physical and social mobility, and social ontology.

Infrastructure

- Characteristics
 - ✓ Technically heterogeneous, geographically dispersed, and institutionally complex without centralized coordination mechanism
- ✓ Challenges
 - Interoperability, scalability, stability, reliability and persistence
 - Seamless intergrations of heterogeneous devices, services, and platforms
 - Location awareness, service availability, physical and social mobility, and social ontology

CASE | Weatherhead School of Management

16

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 16

SERVICES

In ubiquitous computing environments, services need to be configured dynamically with varying lifecycles from many different sources. Thus, personalization and mobility support will be important. In order to support these two aspects, content and medium need to be separated and infrastructure need to provide context awareness to the devices and services.

Services

- Services need to be configured dynamically with varying lifecycles from many different sources
- Challenges
 - ✓ New services
 - ✓ Personalization
 - ✓ Mobility support
 - ✓ Content and medium separation: Management of content and metadata management largely unresolved (despite XML)
 - ✓ Location and time (context) awareness

CASE | Weatherhead School of Management

17

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 17

ORGANIZATIONAL INNOVATION

According to the history of technology-enabled innovations, most such innovations failed not because of the technological breakdowns, but due to the lack of complementary organizational innovations. The integration of both social and physical mobility will foster novel forms of social/technological innovation and demand new ways of organizing. However, often such new forms of organizing are hard to realize. Since the technology is often designed only based on past historical use, it is much more difficult to foresee what the future would look like and build new technology based on such visions. Thus, it is critically important to take co-evolutionary approach in building ubiquitous computing environments through experiments and trial and errors.

Organizational Innovation

- The integration of both social and physical mobility will foster novel forms of social/technological innovation and demand new research approaches
- Driven by both grassroots experiments and novel theoretical and methodological choices
- May demand a radical shift in the focus of research and associated business innovation
- Focus on experimentation, learning from technology trials, novel theorizing
- Designing unforeseen future using technology

CASE | Weatherhead School of Management

18

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 18

ORGANIZATIONAL CHALLENGES

Such experiments and co-evolutionary approaches will eventually help us understand many questions raised in this slide. One of the most critical aspect of organizational challenges is the notion of un-packaging and re-packaging of digital services based on innovative business models. In the past, particular services were tied up with particular physical products. Through digital convergence, however, digital services can be separated from the medium (physical products). Once separated, these contents and media can be re-packaged in different combinations enabling new and novel offering of digital services. Another critical challenge is to figure out how to make money out of this emerging digital service model. In particular, it is very difficult to understand the demand for this type of novel digital services. Also, we have very little understanding of how users consume information rich products. In order to provide offer this type of novel digital services, organizations will have to develop a completely new form of strategic alliances with companies coming from different and remote industries.

Organizational Challenges

- SERVICE CONCEPTS AND STRATEGIES
 - ✓ How to make money?
 - ✓ All information services may need careful rethinking and can be transformed
 - ✓ Design, management and (un-)packaging of digital services based on innovative business models
 - ✓ Understanding user needs in a new information rich environment
 - ✓ Understanding the demand and what drives it is difficult
 - ✓ Evolution and expansion of services based on user learning
 - ✓ Many services based on increasing returns and are community based
 - ✓ New forms of strategic alliances required

CASE | Weatherhead School of Management

19

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 19

INSTITUTIONAL CHALLENGES

Finally, institutional challenges are the most enduring and difficult challenges to overcome. It includes political regulatory challenges. This includes, frequency allocation, access technologies, privacy, security, and address and access types regulations. Often these institutional challenges are embedded into our everyday lives and hard to discern and predict. In order to fully realize the early promises of ubiquitous computing environments, technologists and organizational management alike need to be keenly aware of these institutional challenges.

Institutional Challenges

- **POLITICAL AND REGULATORY CHALLENGES**
 - ✓ diffusion demands coordination on institutional policies: frequencies, access technologies, privacy, security, addresses and access types
 - ✓ Stakes are high nationally and world wide- the race concerning the winners of the next wave of technological transformation will be furious!
- **Institutional barriers**
 - ✓ regulatory regime: innovations in construction and architecture industry

CASE | Weatherhead School of Management

20

Emerging Computing Workshop, NASA Langley
03.19.2003

Figure 20

CATCH THE WAVE!

In order to pursue these emerging opportunities, more than ever, we need interdisciplinary approaches crossing the boundaries between traditional academic disciplines. Furthermore, we will need to combine basic research with application developments, because one leads to the other as discussed earlier. Certainly, the current institutional arrangement at many universities and research laboratories do not make such interdisciplinary research easy.



Figure 21

GRID & AUTONOMIC COMPUTING –THE NEXT EVOLUTION

Nancy Brittle
IBM
Norfolk, VA

Reprinted with the permission of IBM Corporation.

GRID & AUTONOMIC COMPUTING –THE NEXT EVOLUTION

Grid is a key part of IBM's on demand strategy, a powerful vision for the computing enterprise. Our talk today is intended to focus less on the vision and strategy and more on what we have available to help solve your most pressing business and IT problems today.

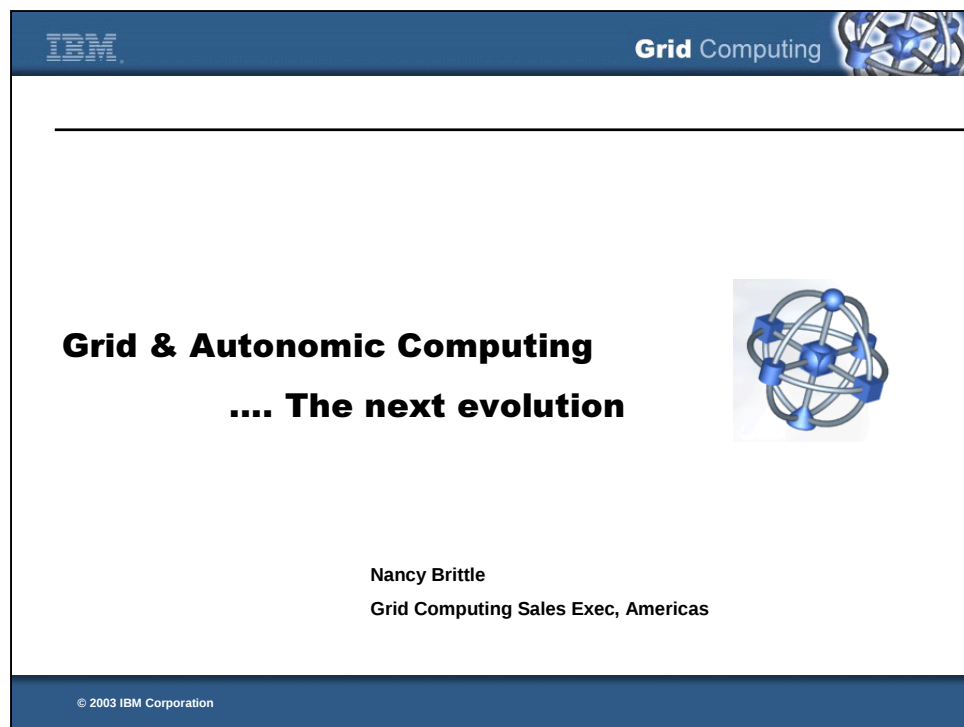


Figure 1

CONTENTS

Grid computing is an emerging technology, but Grid is already delivering real business value to customers today. This first section will focus on the major areas we see that taking place in and show you some examples of what we mean.

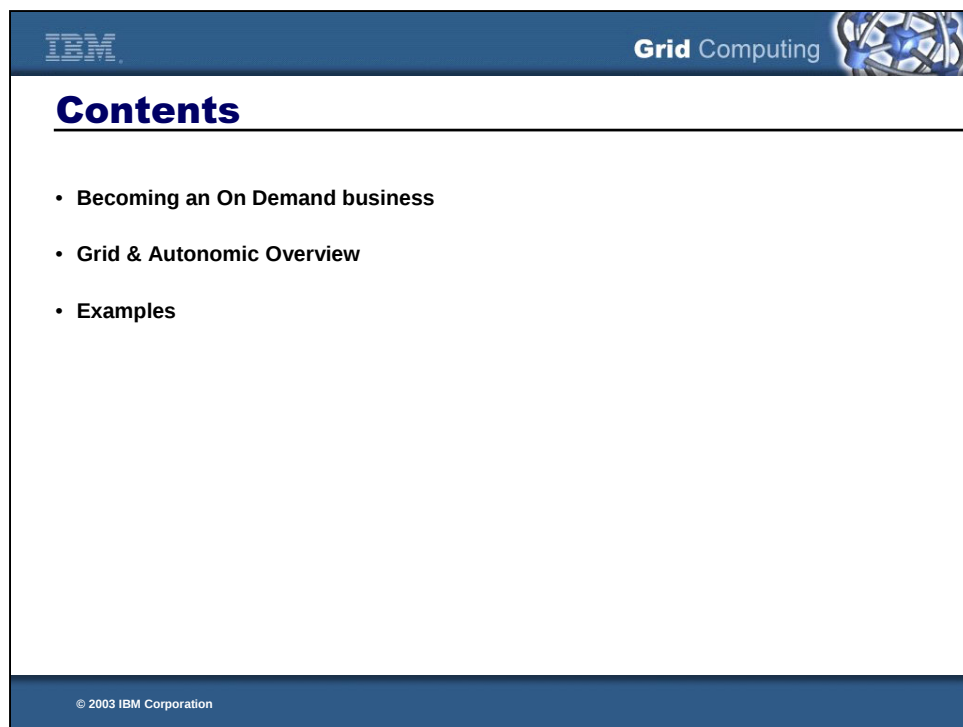


Figure 2

ON DEMAND OPERATING ENVIRONMENT ATTRIBUTES

An On Demand operating environment has these four characteristics:

- 1) Integration...IBM has the middleware products to deliver on this today.
- 2) Open standards...Web services platforms and open architectures to enable rapid deployment and integration of business process applications. IBM embraces open standards across their product brands; this is a key differentiator
- 3) Virtualized: now there's an opportunity to virtualize the entire data center with the emerging technology of grid computing.
- 4) Autonomic: for self-managing systems that include IBM's Tivoli management software and DB2 data base with self-tuning and self-managing features

Grid is a key enabler of the On Demand Operating Environment.

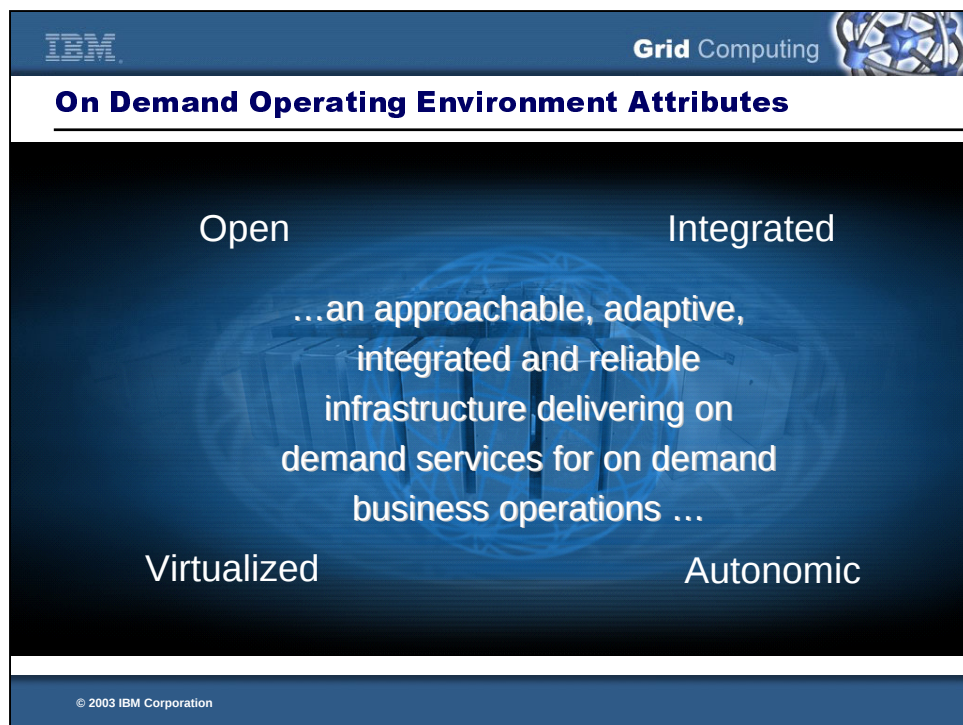


Figure 3

THE ON DEMAND OPERATING ENVIRONMENT

Many new types of input devices and systems need to be integrated. This integration is based on open middleware

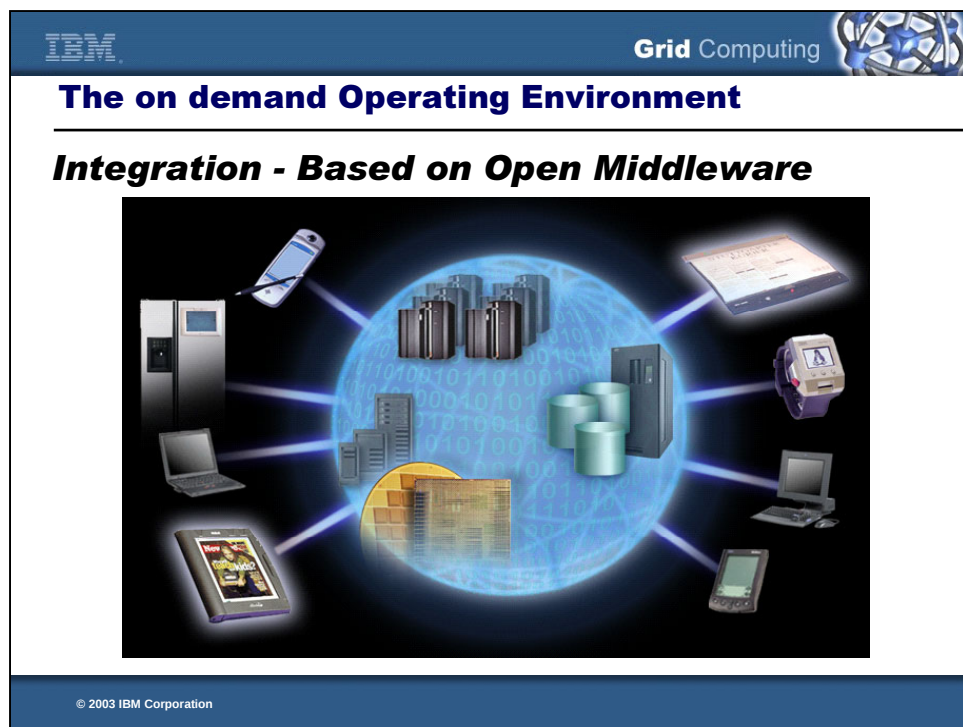


Figure 4

AGENCIES FACE CHANGING MARKET DYNAMICS

The dynamics of the market has changed tremendously over the last few years. Users need technology to adapt to them rather than the users having to adapt to the technology. Unpredictable fluctuations in the market cause uncertainties that business needs to be able to have a flexible environment to handle the dynamic changes. Collaboration among departments, intra-agencies, and globally are becoming more essential than ever before.

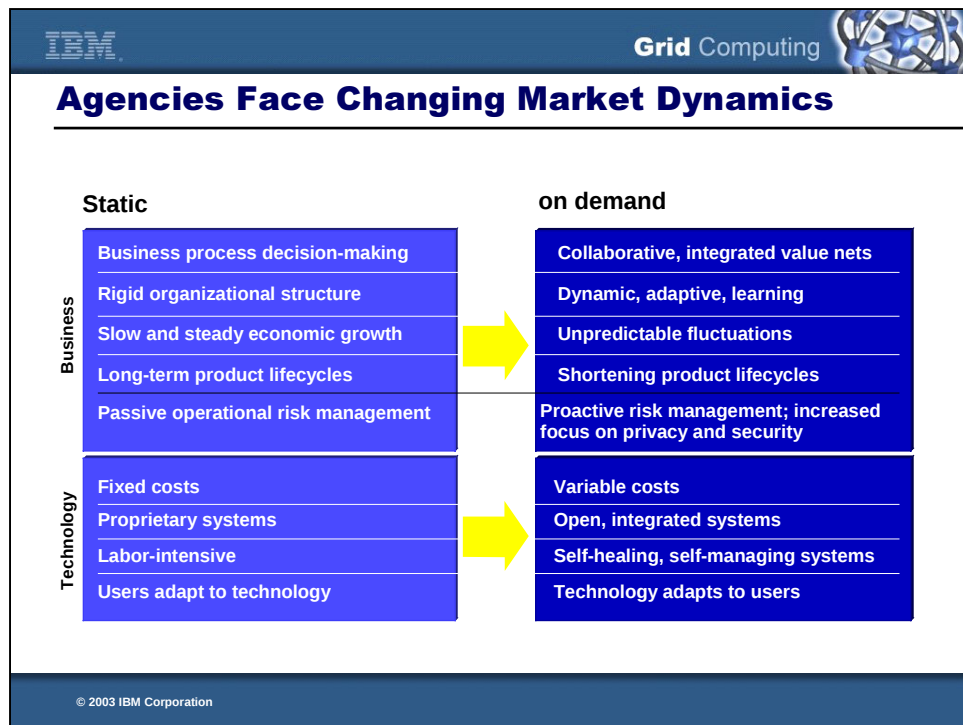


Figure 5

GRID COMPUTING DEFINED

Grids create a virtualized data center: Grids tie together resources across geographical boundaries, organizational boundaries, and system types.

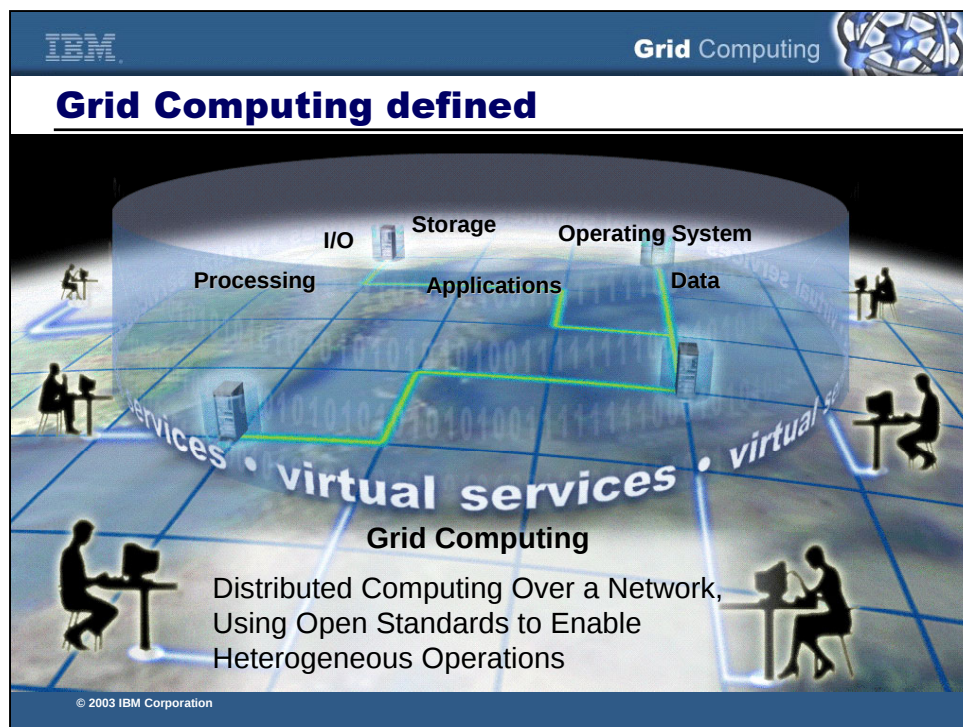


Figure 6

VIRTUALIZATION TECHNOLOGIES TODAY

Virtualizing technologies are not new. We have been using these technologies for some time and are now implementing them across heterogeneous environments.

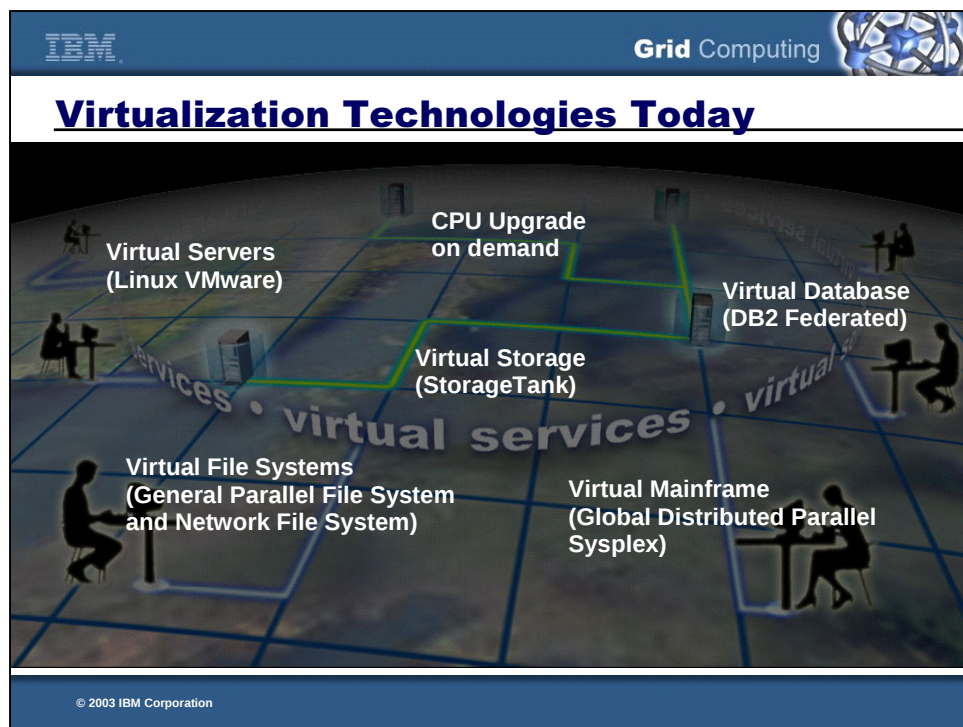


Figure 7

USES OF GRID TECHNOLOGY

There are four models of use with grid technology: Processing grids aggregate the power of heterogeneous servers or desktops to take advantage of unused cycles. Data grids enable disparate data sharing and collaboration across virtual organizations for intelligent decision making using dispersed data and multiple data formats. Resiliency grids enable continuous business operations in case any system in the grid should become enabled due to unplanned or disastrous events occurring. On demand uses the grid architecture and infrastructure to provide a utility for compute resources. Through metering and billing, users will be charged for what the resources they use and charged appropriately...very similar to electricity and water utilities today.

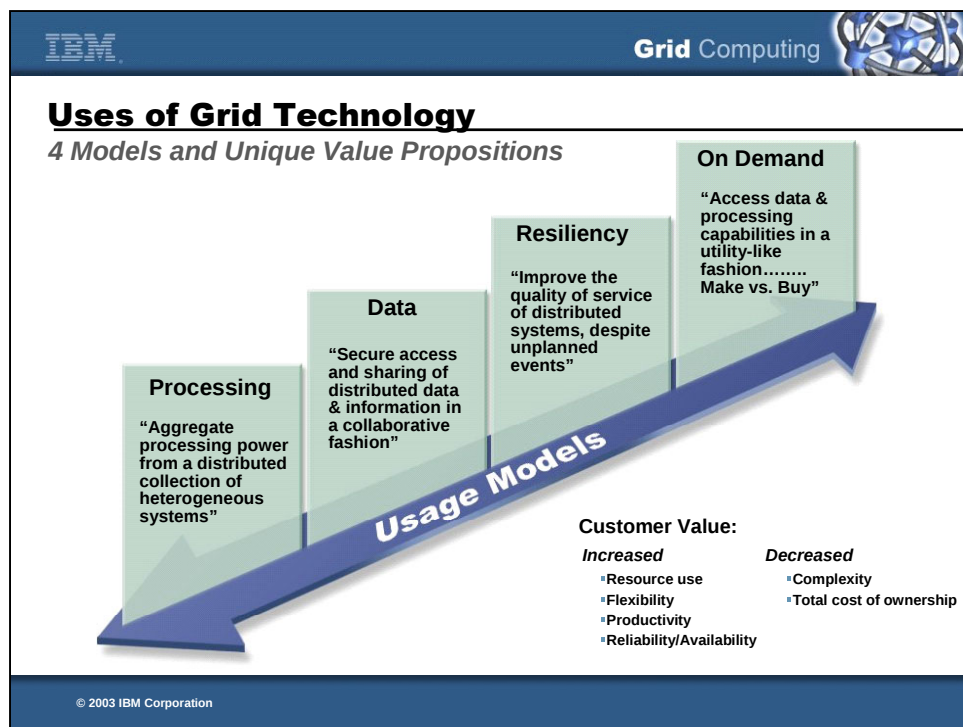


Figure 8

GRID DEPLOYMENT OPTIONS

Customers are deploying grids in many different ways: Intragrids are within a company's firewall to enable inter-department collaboration and sharing of resources. Extragrids connect companies with their suppliers and partners. Intergrid enables collaboration across multiple agencies through the internet. Many researchers in universities begin deploying grids in this manner to enable research data to be shared.

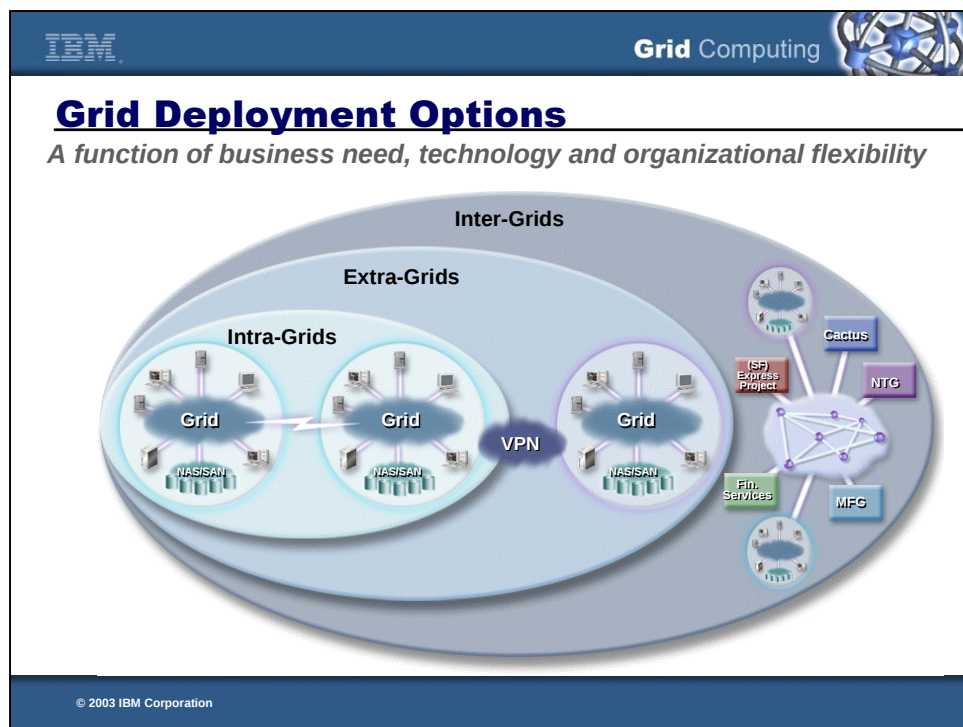


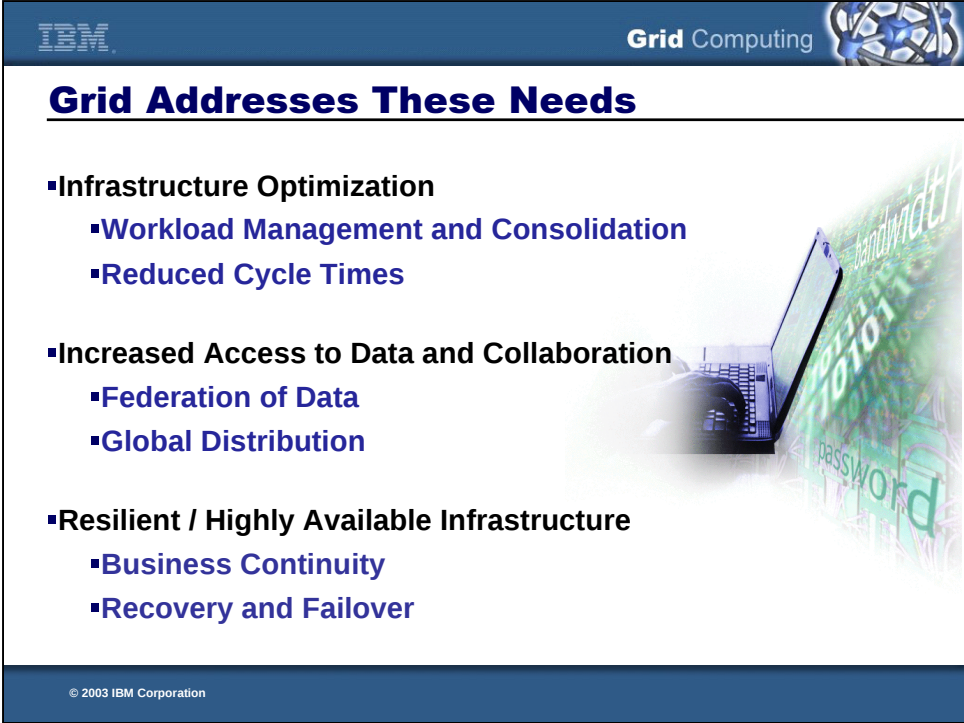
Figure 9

GRID ADDRESSES THESE NEEDS

Grid is delivering real value to businesses today.

Companies are using Grid today in three general areas:

- 1) To improve utilization of computational resources
- 2) To enable collaboration between organizations, and
- 3) To improve the flexibility and resiliency of computing infrastructures.



The slide features the IBM logo and 'Grid Computing' text at the top. The main title is 'Grid Addresses These Needs'. The content is organized into three main bullet points, each with sub-bullets. A graphic of a laptop with floating text like 'bandwidth', 'password', and '101011' is on the right. The footer contains the copyright notice '© 2003 IBM Corporation'.

IBM Grid Computing

Grid Addresses These Needs

- **Infrastructure Optimization**
 - Workload Management and Consolidation
 - Reduced Cycle Times
- **Increased Access to Data and Collaboration**
 - Federation of Data
 - Global Distribution
- **Resilient / Highly Available Infrastructure**
 - Business Continuity
 - Recovery and Failover

© 2003 IBM Corporation

Figure 10

LOW INFRASTRUCTURE UTILIZATION

One of the main drivers for grid computing is the ability for organizations to do more with their currently owned assets. Typically mainframes do a good job at maximizing utilization. Grids can maximize utilization of UNIX and Intel-based systems and can aggregate the collective processing cycles that can work on jobs that were not viable before.

IBM

Grid Computing

Low Infrastructure Utilization

	Peak-hour Utilization	Prime-shift Utilization	24-hour Period Utilization
Mainframes	85-100%	70%	60%
UNIX	50-70%	10-15%	<10%
Intel-based	30%	5-10%	2-5%
Storage	N/A	N/A	52%

Source: IBM Scorpion White Paper: Simplifying the Corporate IT Infrastructure, 2000

© 2003 IBM Corporation

Figure 11

RESILIENT/ HIGHLY AVAILABLE INFRASTRUCTURE

Another significant motivation for employing Grids is the need to reduce the time that it takes to complete a particular computation. Often this need is a critical part of the value proposition for a particular business function.

For example if a particularly compute intensive task must be accomplished in a short span of time opportunities for parallel execution of parts of the calculation can be exploited to complete the job more quickly than if the entire calculation was performed serially. In fact in many cases sequential execution of the problem might take so long as to render the final result unusable.

In another case, advantages may be gained by running a particular computation more often. For if example airline pricing and load management algorithms, which are fairly complex calculations, can be completed more rapidly they can be executed more often allowing the company to respond more rapidly to changing market conditions and better utilize its planes, personnel, and fuel resources.

The animation above shows three jobs that are scheduled to run on three different servers. During the course of running Job 1, the server it is running on has an outage. This might be a Sun server which fails, or it might be a server going down for scheduled maintenance. Its really not important what the reason is. Using Grid middleware, a scheduler can detect that Job 1 did not complete and reschedule that job to run on another available computing resource. This ensures that all critical tasks are completed.

Most corporate computing users don't care where their application runs. They want good performance and they need their data to be secure. Using intelligent scheduling middleware, a company can utilize the most available, appropriate asset to run a given task. It might be the case that during the trading day in New York, the banks data centers in Tokyo or London are idle. The employees there are home sleeping. The middleware can schedule jobs to run in the overseas data centers improving performance of the application and off-loading workload from New York, a win-win scenario.

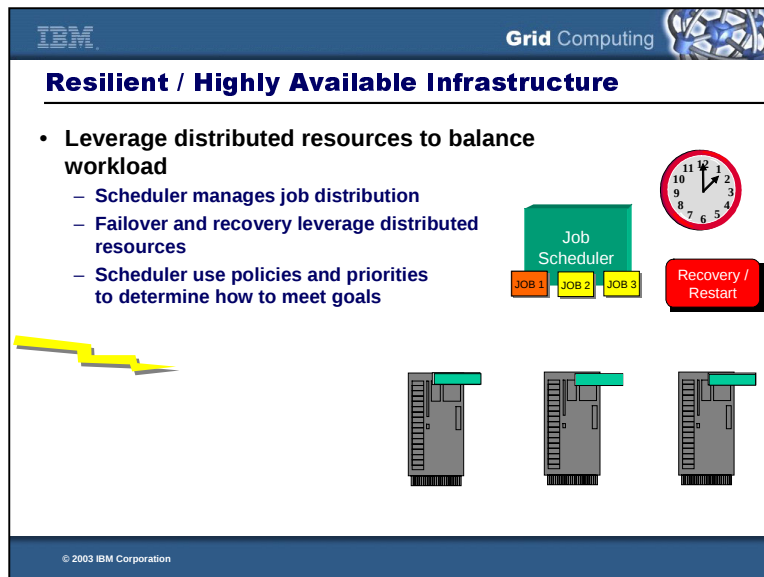


Figure 12

THE VALUE OF OPEN STANDARDS

None of these models will happen without open standards. I think we have seen a pretty clear path over the last 10-15 years of open standards based computing starting all the way back with networking. There were other protocols that came out for networking. SNA, NetBIOS, etc...but people rallied around TCP/IP. It became an open standard approach to be able to take many different computer types and allow them to communicate over a network.

From a communications perspective, we started to see e-Mail packages emerge and we now have standards like SMTP, POP3 and MIME. MIME was a very important standard that allowed different e-mail packages to be able to communicate with one another, and standardized how attachments were handled.

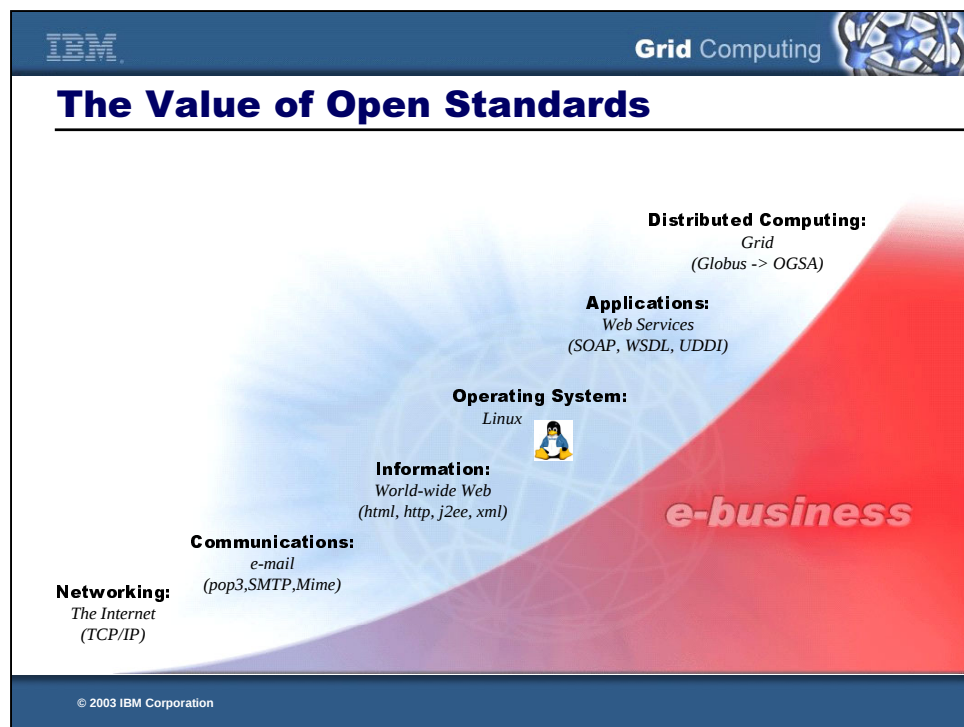


Figure 13

THE VALUE OF OPEN STANDARDS – (CONTD.)

From an information perspective we saw the World Wide Web emerge...again using new protocols and open standards like HTML and HTTP. J2EE has emerged as a standard for the Web infrastructure to communicate with backend transactional systems, your corporate data. Operating systems, the penguin – LINUX has become very popular around the world from a server infrastructure perspective. New feature and functions are being added to LINUX around the world at an incredible pace. The open source community and companies like IBM are participating in this open movement. Because of this LINUX has become a very popular operating system for server environments and becoming almost a defacto-standard as an operating system for servers. From an application perspective, we have the web services standard emerge, focused on SOAP as the transport layer, WSDL as the web services definition layer and UDDI as the directory of web services. Web Services is all about hooking up applications and making application to application communications simpler for developers within an enterprise. Developers can now quickly find web services and assemble them into applications – again, use of open standards driving more value to the business.

And today we are talking about distributed computing - grid computing. And yes again we have a standards body and a process of working with an open standards based community (Global Grid Forum – GGF). IBM is working on the standards for distributed computing with this group. The technology we are all developing is called OGSA, which stands for Open Grid Services Architecture.

If you look at these standards, it is pretty easy to come to the realization that OGSA will be to grid what TCP/IP was to networking, what HTML/HTTP were to the Web. If you want to build a grid of distributed systems & distributed resources you are going to need OGSA on all those platforms and resources within your environment. This is very similar to networking. If you are going to build a network of many different platforms, you need TCP/IP on those platforms for them to communicate.

GRID MIDDLEWARE TODAY

The world of Grid middleware today is very much similar to the early days of networking. These are some examples of Grid middleware. We are working with all of these in various engagements. Each, on its own, has some excellent technical capabilities. But today it is not possible to use the workload scheduling capability from one product, the data management from another and the systems management from another and have confidence in the interoperability of the solution. These are essentially proprietary solutions today as no standard exists yet.

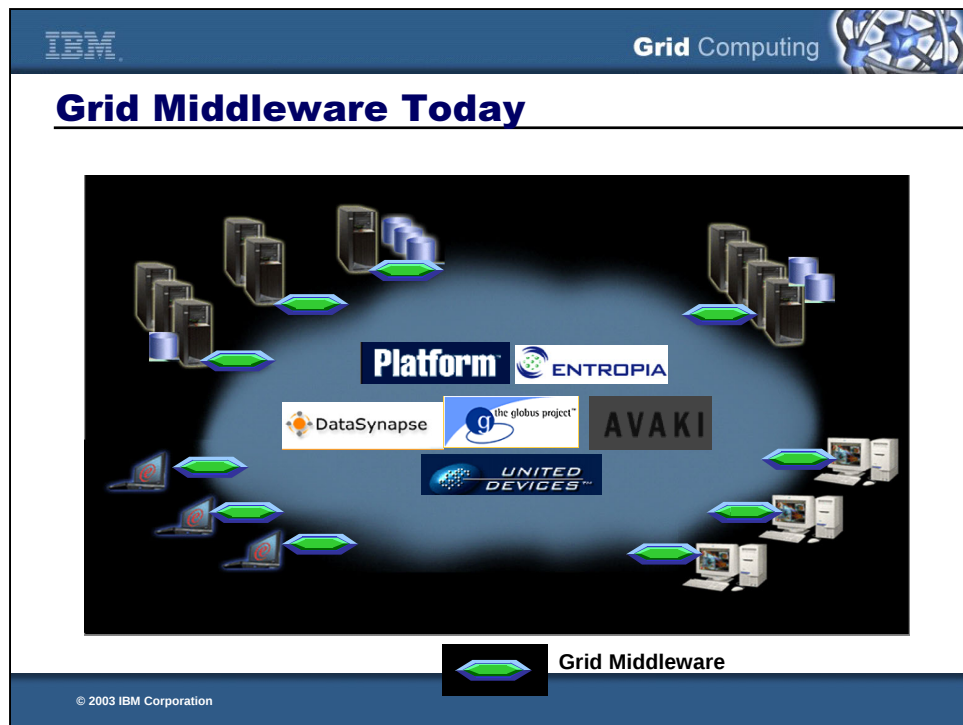


Figure 14

OPEN GRID SERVICES ARCHITECTURE (OGSA)

The Open Grid Services Architecture, or OGSA, will change that. Defined by the Global Grid Forum (or GGF), in which IBM is playing a very active role, the OGSA will be the standard protocol for Grid computing.

The Globus Project is an open source implementation of OGSA (based on the GGF specification) and a toolkit (Globus 3.0) that provides a set of APIs to implement Grid applications.

IBM, and other leading vendors are all sponsors of GGF. As I will show you in a moment, IBM is OGSA enabling all of our related products. Also, all of the middleware providers that I showed you in the previous slide are committed to OGSA and will implement it in their products.

OGSA will be the TCP/IP of Grid computing. It is the common protocol that all computing resources must support to join and interoperate in a Grid. All Grid related middleware will support this standard allowing for interoperability of Grid solutions.

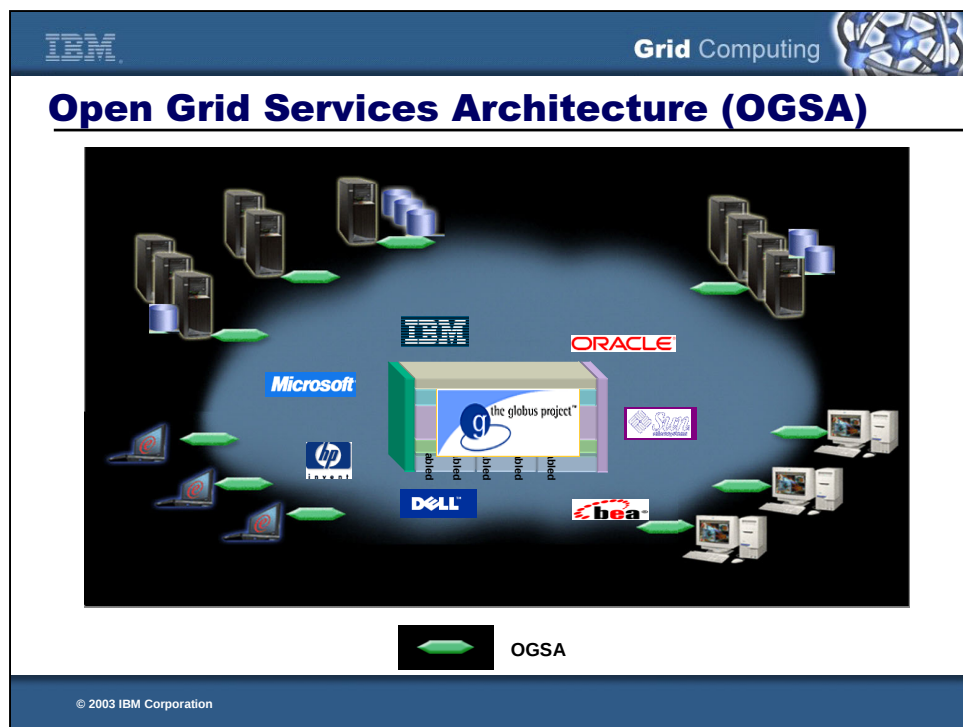


Figure 15

ARCHITECTURE FRAMEWORK OGSA STRUCTURE

So now let's talk about the architecture related to OGSA and Grids.

First, all computing hardware devices that wish to participate in a Grid will be OGSA enabled. This includes servers, storage and I/O devices. And IBM is enabling all of our e-Servers and TotalStorage products.

Next is the general middleware layer where many/most databases, file systems, security services, etc, will be OGSA enabled.

Next is the Web Server engine, the container for the OGSA functionality. There will be many implementations of OGSA carried by open source implementations like JBOSS and by products such as WebSphere, IBM's strategic web engine.

OGSA, the Open Grid Services Architecture -- is being written as J2EE and it will be based on web services. This is a very important point. The developers that are working on the open grid service architecture decided not to recreate the world and they decided to base their work on another standard that is available today -- web services.

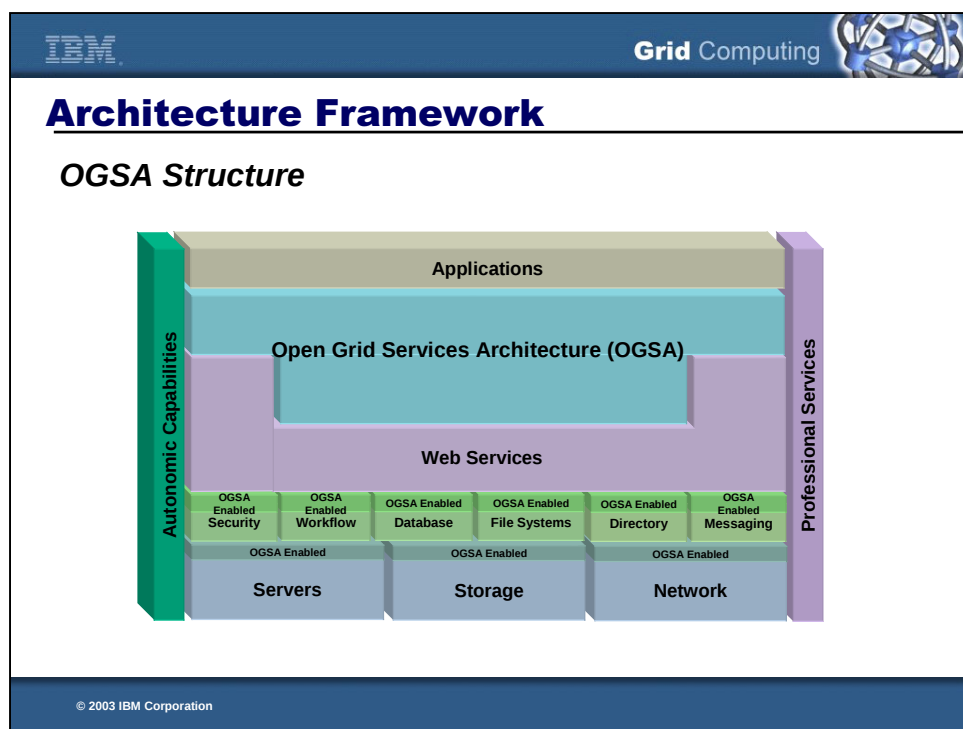


Figure 16

ARCHITECTURE FRAMEWORK OGSA STRUCTURE (CONTD.)

It is a standard that has been driven for the applications developers – to make application development and application integration easier - web services is a perfect way to implement the open standard based grid protocols.

Sitting on top of this stack are the applications that will exploit this functionality.

This function will require autonomic functionality in the infrastructure to keep devices available. We also see a big role for services as the fact is that today grids are built, they are not bought. We believe that our experience and skills participating and helping to build the most significant grids in the world are an important core competency that IBM brings to the table with our clients.

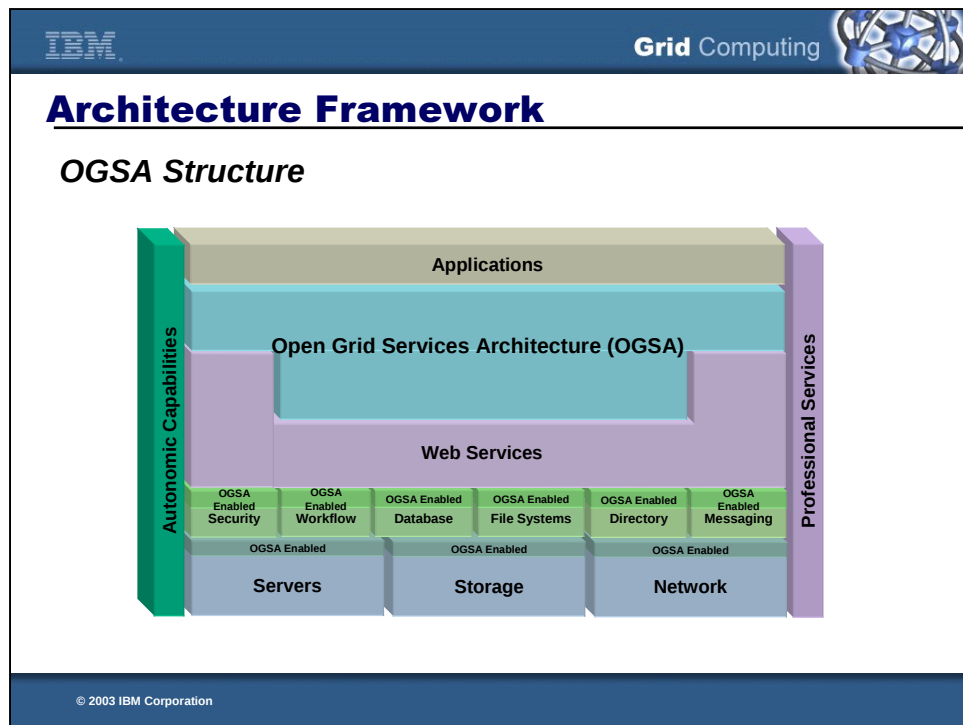


Figure 17

ARCHITECTURE FRAMEWORK PRODUCTS AND SERVICES FOR GRIDS

Here you see the total stack populated with the types of products we expect to see running in the Grid world in the future. Please note that there are no commitments from these application ISVs at this time, but I want you to understand that Grid is not just about HPC applications. These capabilities will open up important functions for mainstream business applications as well.

Some of IBM's key Grid partners are in this picture - Platform, Avaki, DataSynapse, Entropia, United Devices, and the Globus Project. Today they have software - grid middleware that allows customers to build grids. In the future they will be recasting their products to work on top of OGSA. What they have all realized is that we don't need 7 or 10 proprietary ways of building grids in the world – just like we did not need 7 to 10 ways to do network in the world. We need one open standard way that all customers can depend on, a standard that allows Grid ISVs to provide higher level grid services. Customers and application ISVs can be assured that there is one open standard way of building and deploying these services, this is what OGSA is all about.

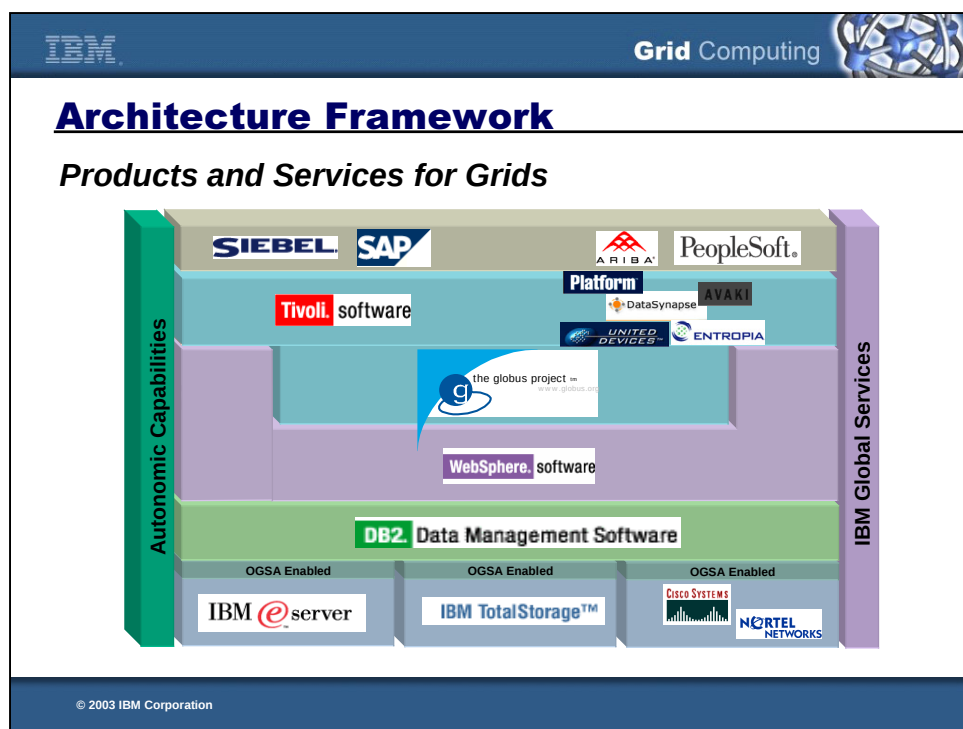


Figure 18

ARCHITECTURE FRAMEWORK PRODUCTS AND SERVICES FOR GRIDS – (CONTD.)

As you see in the middle of this picture, OGSA will need an application server and Web services engine. At IBM we are very excited about this architecture because we feel we have developed and continue to improve on the best web services engine in the world - WebSphere. We expect to make many enhancements to our WebSphere product with respect to web services. We expect to be able to run web services better than anyone in the industry on multiple platforms. We intend to provide the highest level of resiliency in the industry, and the highest Quality of Service for web services. We have been told by our customers that supporting multiple platforms, resiliency, QoS, and open standard are some of the most important things they want in their IT infrastructure. We will deliver Grids through WebSphere. Tivoli products will be enhanced and focused on grid deployment and management, and our storage and database products are being enhanced to support Grids.

As these architecture shows OGSA will be used an this open standard based protocol that will support multiple servers, operating systems, storage & data systems in a very resilient fashion. This is the architecture of grid and of future IT environments.

AUTONOMIC COMPUTING

Why is Autonomic Computing important in a Grid environment? To net it out, it means that systems are self-configuring, self-healing, self-optimizing and self-protecting; it means that systems do the work, freeing IT professionals to focus on other critical business needs.

- Autonomic computing systems are self-configuring, self-healing, self-optimizing and self-protecting.
- Self-configuring systems increase IT responsiveness/agility
- Self-healing systems improve business resiliency
- Self-optimizing systems improve operational efficiency
- Self-protecting systems help secure information and resources

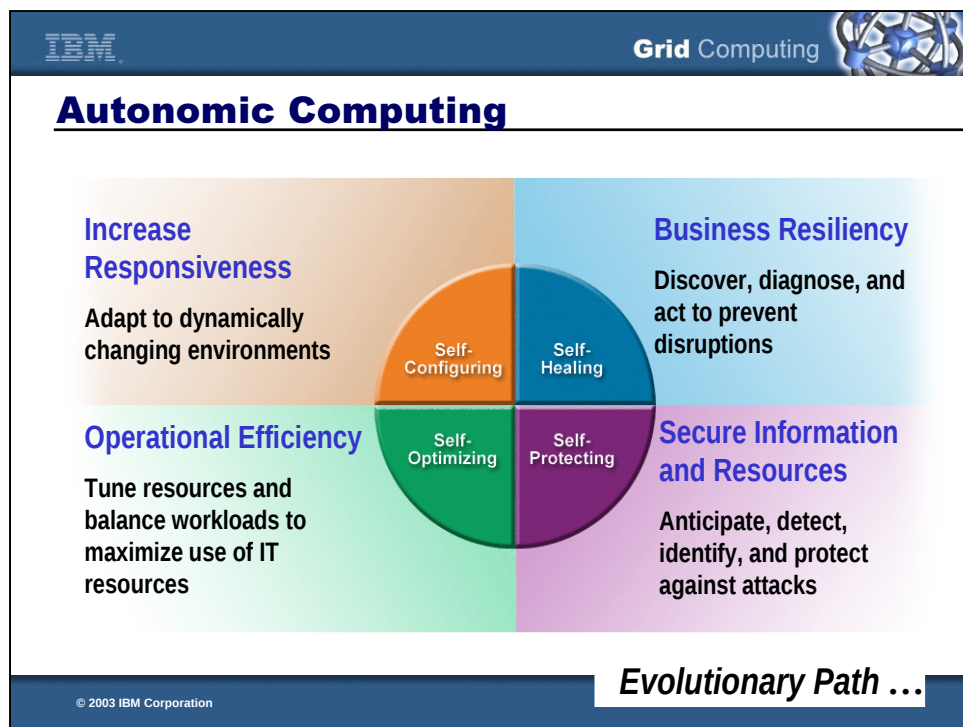


Figure 19

AUTONOMIC EXAMPLES

How does Autonomic computing fit into Grid?

Autonomic capabilities are found in Grids today as they are already available in many IBM products today. Including...

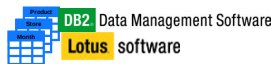
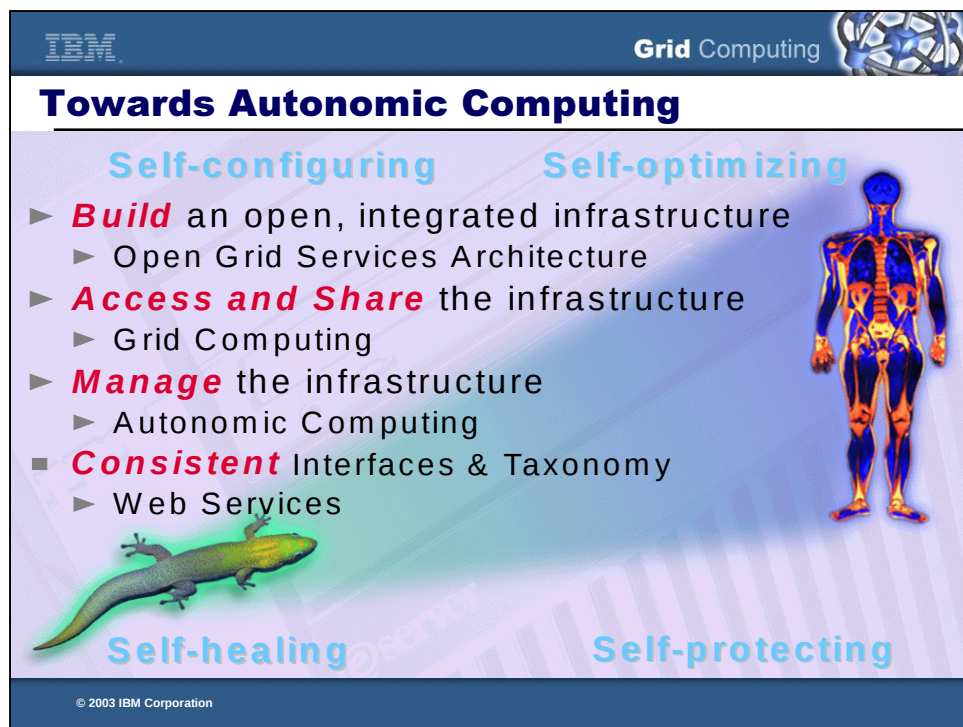
 Grid Computing 		
Autonomic Examples		
Systems Management		✓ Access / Identity Managers ✓ Storage Resource Manager ✓ Service Level Advisor
Client		✓ ImageUltra ✓ Rapid Restore PC ✓ Embedded Security Subsystem
Application		✓ Prioritization of User Transactions ✓ Custom Advisors ✓ Problem Analysis and Recovery
Database & Collaboration		✓ DB2 Query Patroller ✓ Tivoli Analyzer for Domino
Servers		✓ Dynamic Partitioning ✓ IBM Director ✓ BladeCenter
Storage		✓ Intelligent cache configuration ✓ Predictive Failure Analysis ✓ Dynamic volume expansion
© 2003 IBM Corporation		

Figure 20

TOWARDS AUTONOMIC COMPUTING

OGSA enables an open integrated infrastructure to be built.
Grid computing means accessing and sharing the infrastructure
Autonomic helps in managing the infrastructure
Web services provide consistent interfaces and taxonomy



The slide features the IBM logo in the top left and 'Grid Computing' with a molecular structure icon in the top right. The title 'Towards Autonomic Computing' is centered. The main content is a bulleted list with four primary items, each with a sub-item. The items are: 'Build' (with sub-item 'Open Grid Services Architecture'), 'Access and Share' (with sub-item 'Grid Computing'), 'Manage' (with sub-item 'Autonomic Computing'), and 'Consistent' (with sub-item 'Web Services'). The slide is decorated with a blue and yellow human figure on the right, a green gecko on the bottom left, and the words 'Self-healing' and 'Self-protecting' at the bottom. A copyright notice '© 2003 IBM Corporation' is at the very bottom.

IBM Grid Computing

Towards Autonomic Computing

- **Build** an open, integrated infrastructure
 - Open Grid Services Architecture
- **Access and Share** the infrastructure
 - Grid Computing
- **Manage** the infrastructure
 - Autonomic Computing
- **Consistent** Interfaces & Taxonomy
 - Web Services

Self-healing Self-protecting

© 2003 IBM Corporation

Figure 21

GRID ADOPTION CURVE

The early adopters of grid computing began in universities and research scientists who needed more and more compute power and couldn't afford the cost of the supercomputers to do their work. Aggregating the capacity of multiple computers provided an answer to their problems. We continue to see every industry building grids today and learning about the business value that grids bring. While high performance, numeric intensive environments were the early adopter application drivers, commercial applications are returning significant ROI for businesses today and the trend will continue to increase.

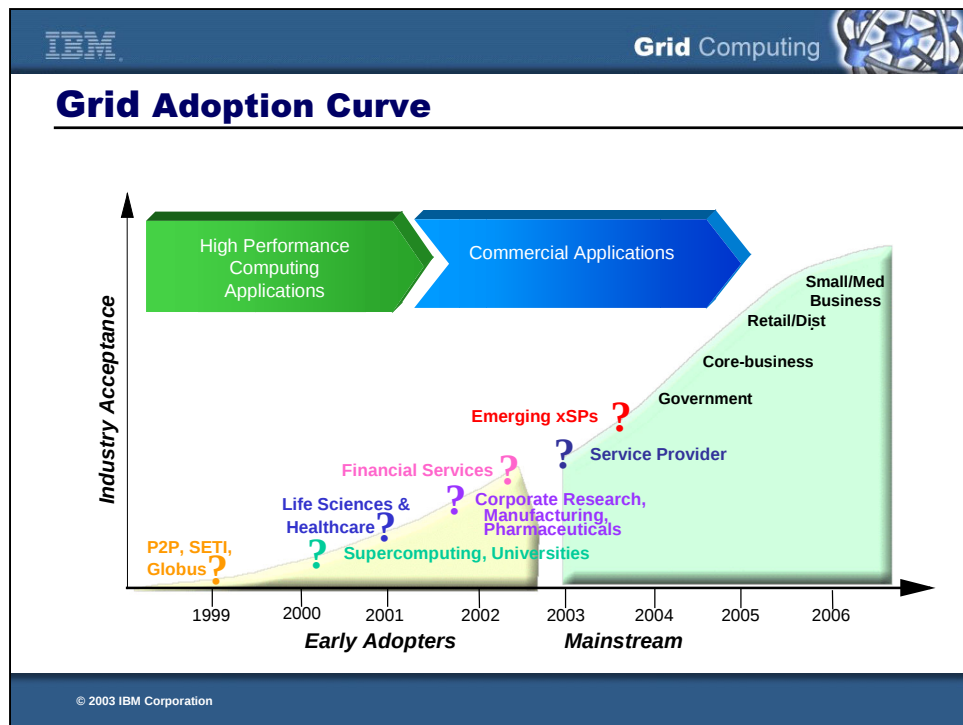


Figure 22

INDUSTRY APPLICATIONS

Grids are being built in every industry today. Some of the key applications within these industries are shown here.

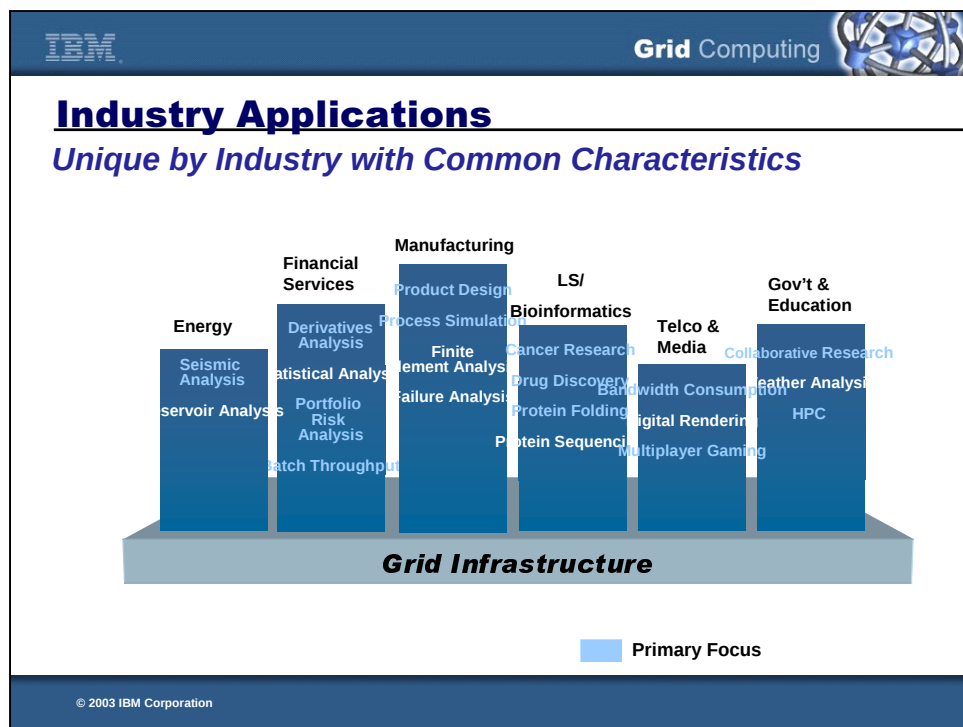


Figure 23

SCIENTIFIC GRID EXAMPLES

Some examples of grids are shown but more information on each can be found on our website at www.ibm.com/grid



Scientific Grid Examples

UK Research Grid Collaborative, scientific research 	The TeraGrid Most powerful, heterogeneous Grid 
University of Pennsylvania - National Digital Mammographic Archive - Medical and diagnostic content  	Indiana Virtual Machine - Purdue and Indiana University 1.4 teraflops - Ultra-large calculations - Simulating homeland security 

© 2003 IBM Corporation

Figure 24

R&D GRID: THE TERAGRID

An example of a large grid implementation is the TeraGrid at NCSA, San Diego Supercomputing Center, Argonne National Lab, and CalTech.



R&D Grid: The TeraGrid

Heterogeneous Systems

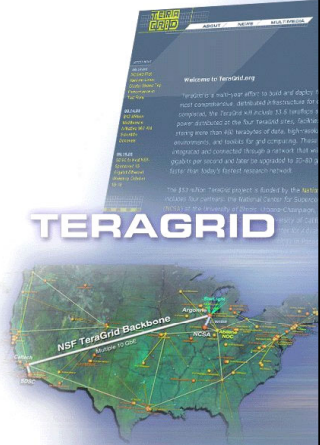
- National Center for Supercomputing Applications
- San Diego Supercomputing Center
- Argonne National Laboratory
- California Institute of Technology

13.6 trillion floating point operations per second

600 terabytes of data

40 gigabits per second

Accessible to thousands of scientists working on advanced research



TERAGRID

© 2003 IBM Corporation

Figure 25

COMMERCIAL EXAMPLES

More examples showing commercial application use. IBM is 'eating our own cooking' by using grids in many areas of our business.



Commercial Examples

John Deere ▪ IBM DB2 DataJoiner ▪ VLDB Award Largest 'federated' database ▪ 2.5 Billion Records ▪ DB2, Oracle, MS SQL 	Financial Services ▪ Linux based Grid ▪ Blade, Grid, Utility technology ▪ Credit Derivatives, Risk Management, Interest Rate Derivative Analysis 
IBM Chip and Server Design ▪ 120TB of WW File Sharing & Collaboration ▪ Thousands of Unix servers distributing workload 	Aventis ▪ IBM DiscoveryLink ▪ Drug discovery collaboration ▪ Diverse, highly distributed 

© 2003 IBM Corporation

Figure 26

IBM

IBM has all of our research centers around the world on an intragrid as well as using grids in our manufacturing plants, benchmarking centers, and design centers.



IBM

Enterprise Optimization

Design Centers for e-business on demand

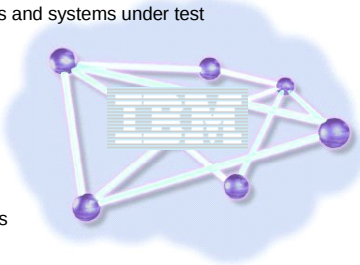
- Virtualize resources in Design Centers to leverage worldwide capabilities
- Mimics customers' distributed environments
- Enables On Demand proof-of-concepts which are integrated, open, autonomic and virtualized

eServer Benchmarking Mercury LoadRunner

- Improved scalability testing by utilizing the full complement of resources in the Benchmarking Centers
- Increased flexibility by breaking co-location of test engines and systems under test

eServer Benchmarking Production Grid and Solutions Marketing Grid

- Enabling On Demand organizations by providing a single interface to access resources
- Reduced workload of staff equates to cost savings
- Increased customer satisfaction by providing resource information and allowing advance reservation of resources

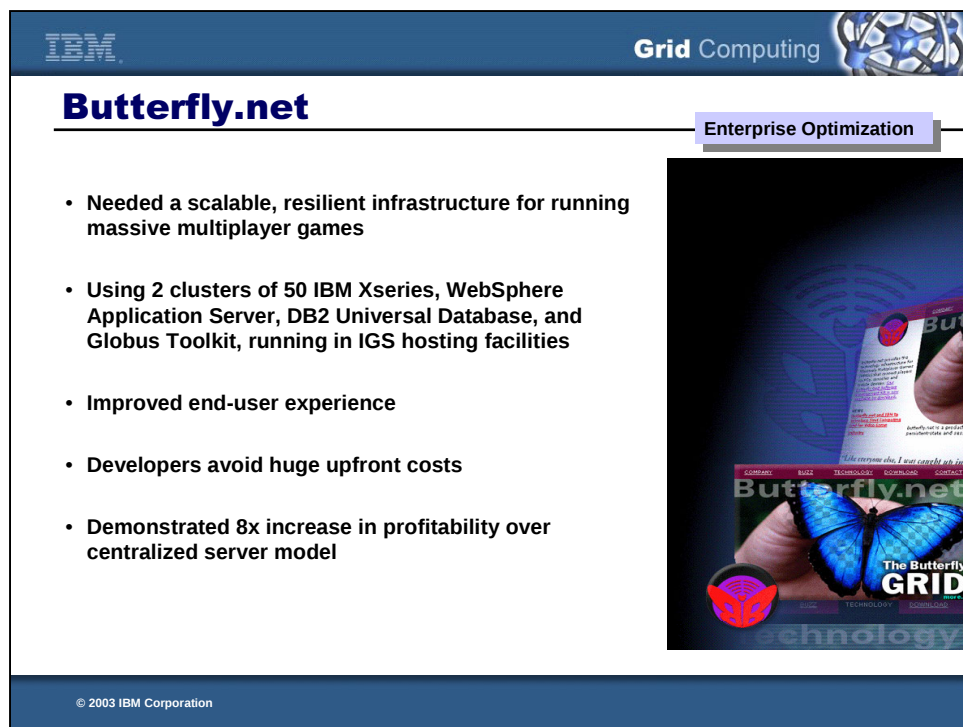


© 2003 IBM Corporation

Figure 27

BUTTERFLY.NET

Butterfly.net is a commercial application of a grid showing significant business value for a company in the video gaming business.



The slide features the IBM logo and 'Grid Computing' header. The main title is 'Butterfly.net' with a sub-header 'Enterprise Optimization'. A bulleted list on the left details the infrastructure and outcomes. On the right, an image shows a hand holding a device displaying the Butterfly.net website, which includes a butterfly logo and the text 'The Butterfly GRID technology'.

IBM Grid Computing

Butterfly.net

Enterprise Optimization

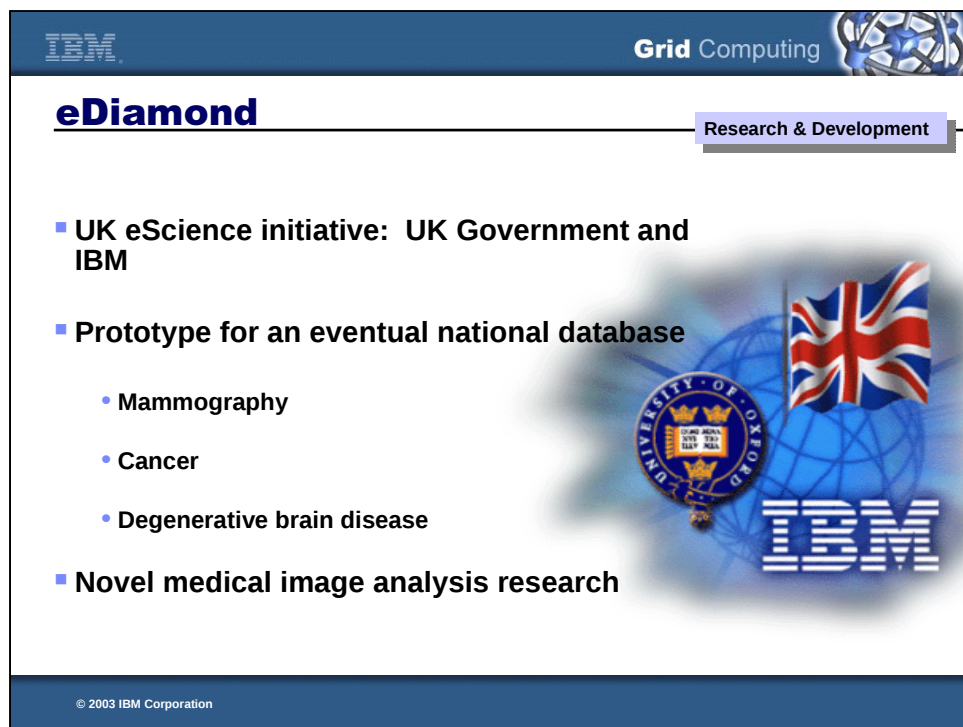
- Needed a scalable, resilient infrastructure for running massive multiplayer games
- Using 2 clusters of 50 IBM Xseries, WebSphere Application Server, DB2 Universal Database, and Globus Toolkit, running in IGS hosting facilities
- Improved end-user experience
- Developers avoid huge upfront costs
- Demonstrated 8x increase in profitability over centralized server model

© 2003 IBM Corporation

Figure 28

eDIAMOND

The UK eScience initiative is designed to tackle issues dealing with cancer research and shortening the time to find cures.



The slide features a dark blue header with the IBM logo on the left and 'Grid Computing' with a molecular structure icon on the right. The main title 'eDiamond' is in large blue font, with a 'Research & Development' tag to its right. The content area lists four bullet points: 'UK eScience initiative: UK Government and IBM', 'Prototype for an eventual national database' (which includes sub-bullets for 'Mammography', 'Cancer', and 'Degenerative brain disease'), and 'Novel medical image analysis research'. On the right side of the content area is a graphic with the University of Oxford crest, the Union Jack flag, and the IBM logo. The footer contains the copyright notice '© 2003 IBM Corporation'.

IBM Grid Computing

eDiamond

Research & Development

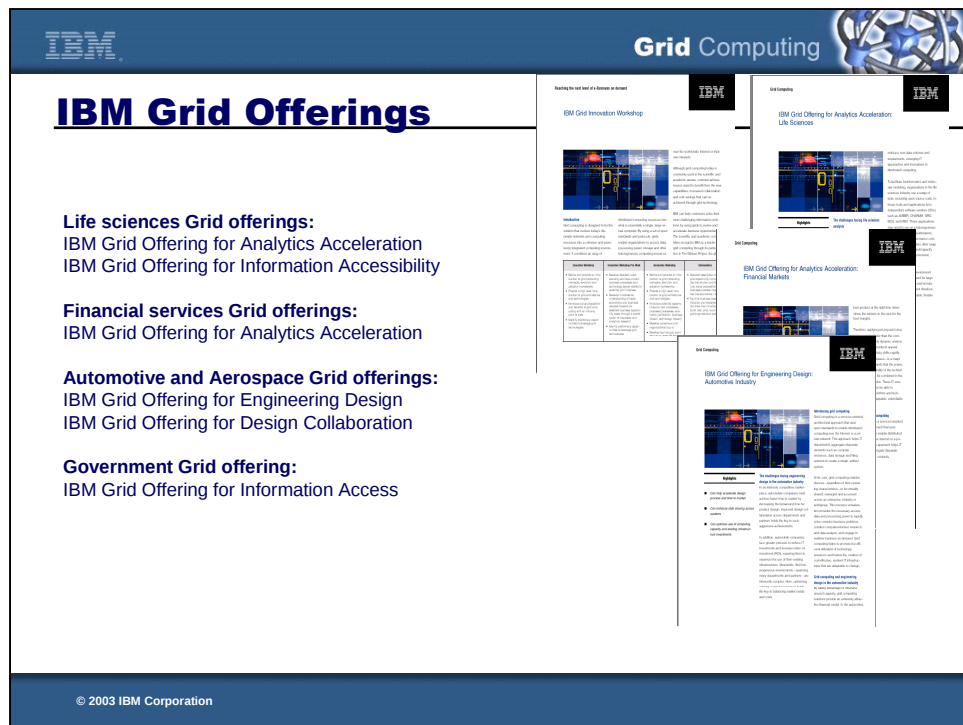
- UK eScience initiative: UK Government and IBM
- Prototype for an eventual national database
 - Mammography
 - Cancer
 - Degenerative brain disease
- Novel medical image analysis research

© 2003 IBM Corporation

Figure 29

IBM GRID OFFERINGS

IBM has announced several Grid Offerings. Targeted toward industry segments, the offerings provide options in middleware, hardware accelerators, etc to meet the needs of each customer's business problems.



The slide is titled "IBM Grid Offerings" and features the IBM logo in the top left corner. The main content is organized into a grid of sections, each with a title, a brief description, and a small graphic. The sections are:

- Life sciences Grid offerings:**
 - IBM Grid Offering for Analytics Acceleration
 - IBM Grid Offering for Information Accessibility
- Financial services Grid offerings:**
 - IBM Grid Offering for Analytics Acceleration
- Automotive and Aerospace Grid offerings:**
 - IBM Grid Offering for Engineering Design
 - IBM Grid Offering for Design Collaboration
- Government Grid offering:**
 - IBM Grid Offering for Information Access

On the right side of the slide, there are four smaller sections, each with a title and a graphic:

- IBM Grid Innovation Workshop**
- IBM Grid Offering for Analytics Acceleration: Life Sciences**
- IBM Grid Offering for Analytics Acceleration: Financial Markets**
- IBM Grid Offering for Engineering Design: Automotive Industry**

At the bottom of the slide, there is a copyright notice: "© 2003 IBM Corporation".

Figure 30

IBM GRID FOCUS AREAS

Grid is most often being implemented in these 5 areas:

- Government Development
- Enterprise Optimization
- R&D
- Engineering & Design
- and Business Analytics.

IBM has 10 offerings available today in these 5 focus areas for select industries:

- Life Sciences
- Financial Services
- Aerospace
- Automotive
- Government

Grid Computing

IBM Grid Focus Areas

Government Development	Enterprise Optimization	Research & Development	Engineering & Design	Business Analytics
Create large-scale IT infrastructures to drive economic development and/or enable new collaborative government services	Optimize computing and data assets to improve utilization, efficiency and business continuity	Accelerate and enhance the R&D process by enabling the sharing data and computing power seamlessly for research intensive applications	Share data and computing power, for computing intensive engineering and scientific applications, to accelerate product design	Enable faster and more comprehensive business planning and analysis through the sharing of data and computing power

IBM has a systematic offering in each focus area.

© 2003 IBM Corporation

Figure 31

GRIDS DELIVER BUSINESS VALUE

Grids can deliver real business value today. Even though the standards continue to evolve and the journey to on demand computing may take a few years to mature, building the infrastructure today is key.

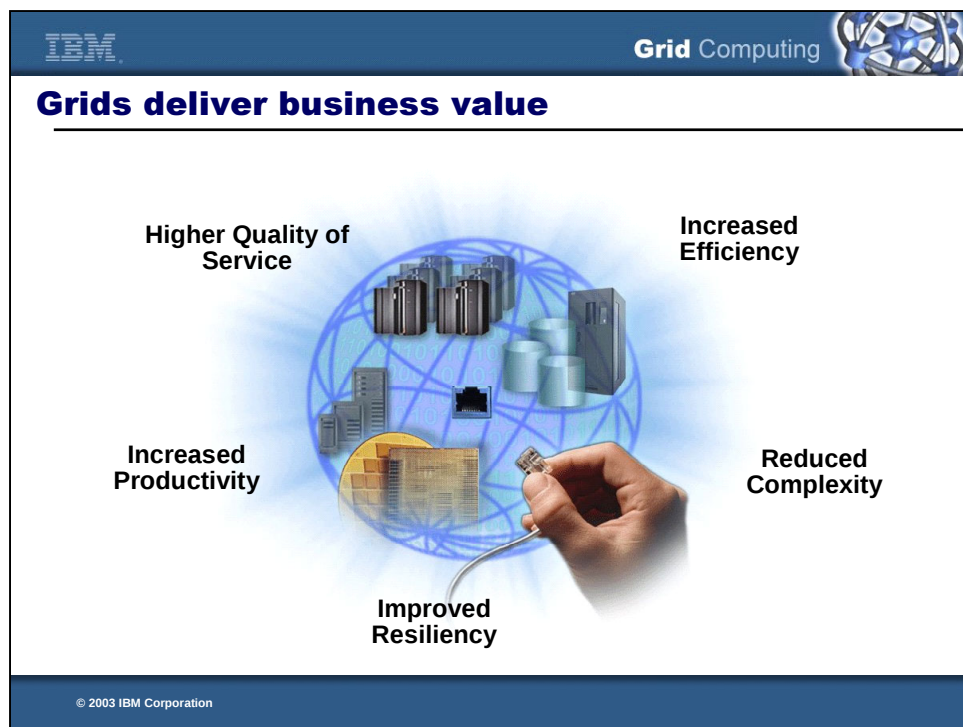


Figure 32

GRID COMPUTING – CHANGING THE IT INFRASTRUCTURE

Steve Salkeld
Platform Computing
Brampton, Canada



Platform™

Platform™

Founded in 1992

Over Ten Years experience Distributed Computing & Grid

400 + employees

50% average growth in the past 4 years

Profitable every quarter since inception

1600+ customers

55% of Top 20 Pharmaceutical Companies

60% of the Top 50 global 500 Companies

90% of top 20 Automotive Companies

80% of Top 20 Industrial manufacturing Companies

Working with world's largest global Financial Institutions

Significant presence in largest Government and Research Centers

Figure 1



Platform™

Agenda

- IT issues – is GRID a solution?
- Platform Products
- Case Studies

Figure 2

TODAY'S BUSINESS CLIMATE

Adopting to the changing market place demands have left businesses challenged to rethink their approach to IT. The spending approach that was identified with the dot-com era or pre-recession times is no more. The investment decisions must be made carefully and in a real context of the demands of the business. It is no longer acceptable for IT to manage boxes – 99.99% uptime for networks is useless if the database the application depends on is down. Providing the complete view of the servers, applications, web services and data along with the people who are using them is increasingly the minimum stakes for IT. The under-current from the past few years have lead to the financial arm of the corporation expecting more then “hand-waving” ROI. It must be clear and linked to the business. The one constant in today's marketplace is change. It falls on IT to adapt to the changing demands of business units, economic trends and strategic decisions made by the corporation.

Platform

Today's Business Climate



- Accomplish more with less
- Transform IT from a “Systems Provider” to a “Services Provider”
- Prove “Hard ROI” for IT investment decisions
- Quickly adapt IT services to changes in business or market demands

Figure 3

IT DRIVERS

These are the IT drivers, in fact the corporate drivers that continuously come up with the customers we have spoken with are as follows: Planning – IT investment decisions are often made in a vacuum with little context of past performance. For example, the switch from Solaris to Linux may save upfront capital expenditure, but how will the applications perform? Under normal load? Under production load? During peak demand? What about provisioning for peak loads? Is it the right methodology?

Server Consolidation – is a practice driven by merged operations, the demand to simplify the management of distributed operations. Looking at the Life Sciences market, the global merging of corporate entities means overlapping IT and business units. The ability to successfully implement a consolidation program is predicated on understanding where to make the best, most effective changes.

Business Continuity – both availability and service levels underscore the need for the consistent delivery of IT services. By understanding the full breadth of services in terms of the time of day, key service windows, and holistic view of all the components of the services.

ROI – By providing clear ROI, in terms of cost, productivity and value during the key windows of corporate performance – market open, product data management load time, B2B uptime will provide a real-world accountability for these services.

Platform

IT Drivers

-  **Planning**
 - Confident purchasing decisions
 - Allocation aligned with need
-  **Server Consolidation**
 - Global historical view of operational centers
 - Improved program success
-  **Business Continuity**
 - Confident delivery of complex services
 - Strengthen the processes for operational availability
- ROI**
 - Strengthen enterprise management
 - Assured accountability

Figure 4

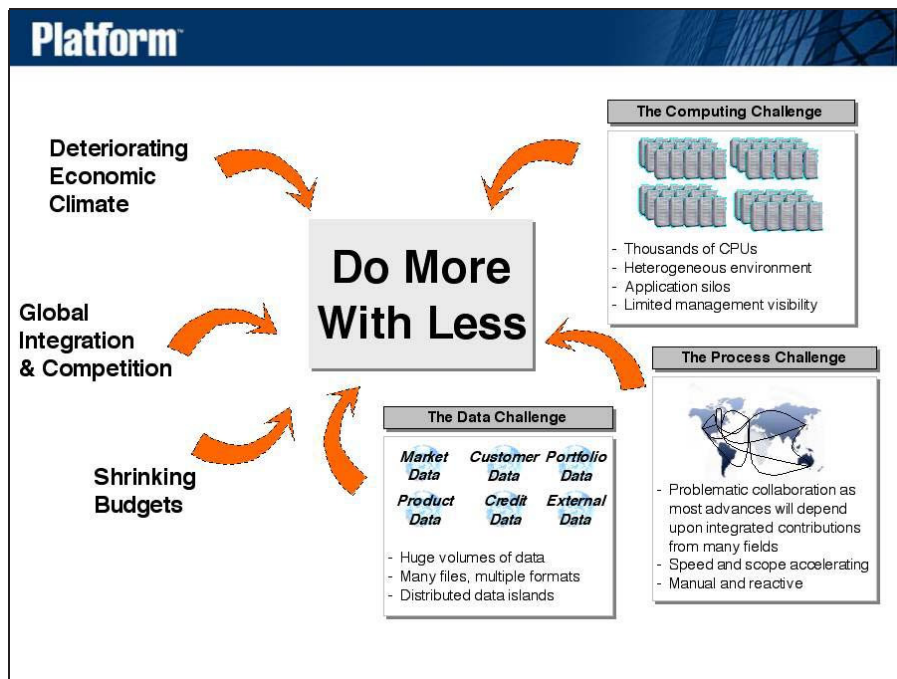


Figure 5

Platform

A large auto manufacturer questions

Why hasn't the goal been achieved?

...because IT is ALWAYS moving, ALWAYS changing!

- Autonomous Organizational Silos
 - prohibits vendor integration
- Highly complex products that require specialized resources
- Present e-Management mechanisms have proven to be incapable of solving the problem
- Reactive, rather than pre-emptive and proactive
- Policy & procedure red tape
- Present manual documentation process can not keep up with the rate of change

7

Figure 6

WHY ARE WE HERE?

GM has architected a new system that will enable GM to achieve its enterprise management goal. The NGM team has proven that it can deliver it!

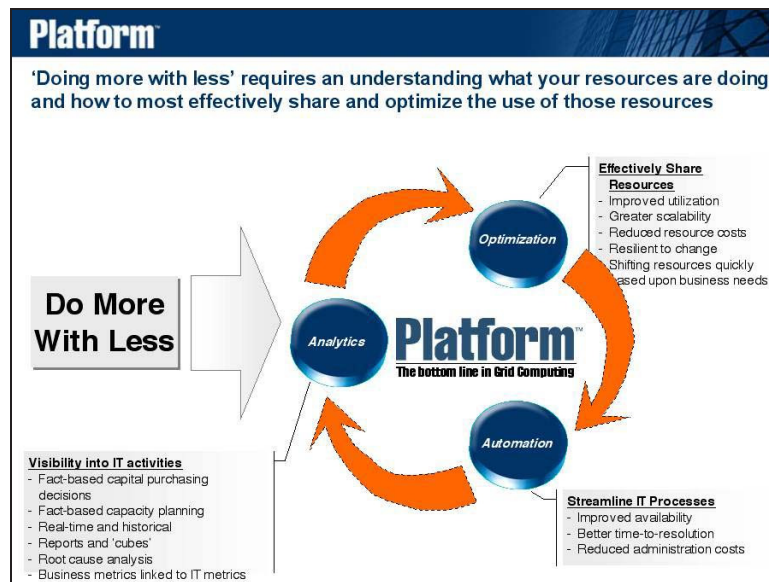


Figure 7

Platform™

A large auto manufacturer says

“ In order to convert [enterprise management] data into information, and transform information into knowledge, one must first truly know the circuit being managed. Automating the documentation of the physical architecture, logical architecture, and detailed bill-of-materials is the first step to understanding where you're at, so that you can begin to get where you want to go with managing distributed computing assets”

9

Figure 8



Figure 9



Figure 10

GRID EVOLUTIONARY NOT REVOLUTIONARY

When you look at deploying Enterprise Grid capability, we believe is Grid is **evolutionary NOT revolutionary** – you need to take steps to implement the technology:

1. First, you need to connect the assets. This enables you to see what you have, where you can first deploy grid technology, and allows for distributed processing.
2. You can't manage until you measure, so the next step is to understand the drivers – who are using what, at what time of day etc. Understand the dynamics of how your resources are being utilized. This is a measurement exercise.
3. From that understanding, you are now in a position to manage and provision the resources more intelligently. This is done with business policy engines, self-healing/HA software and intelligent resource provisioning.
4. Once we are managing, we are now in a better position to now tune the infrastructure based on user demand. – start getting smart about your work in the context of business priorities. This is driven by tight integration with the user applications.

As you build up the Grid framework you will notice that the business silos disappear – as we connect and optimize, you are using everything you have in a much more effective way – increasing collaboration, utilization and delivering better ROI to users and delivering a much better return on the IT asset.

GRID FABRIC: DC INFRASTRUCTURE FOR CLUSTERING AND GRID (COMES WITH ALL OF OUR PRODUCTS)

- *Resource Agent*: Functions to gather information and operate on any type of resources, agent framework for new agent development and plug-in
- *Communication Backbone*: scalable, reliable, efficient and extensible infrastructure to collect resource data and execute actions across grid
- *Distributed Task Execution*: facilities to perform user jobs and management tasks on any devices and resources across grid
- *Central Management*: creation of a “virtual mainframe” infrastructure

Performance Management: Measure and analyze system and application performance against business requirements:

- *System and Application Metrics*: Key performance indices related to business
- *Grid Reporting*: performance reporting, resource accounting and charge back
- *Grid Planning*: bottleneck identification, capacity planning, workload policies
- *Management Portal*: for transparent, secure access to grid management info

Service Management: Manage resource supplies to deliver to the most critical work activities:

- *Self-Healing Management*: Automation of administrative tasks to keep systems and applications in working order and reduce admin costs
- *Service Provisioning*: Dynamic allocation and aggregation of resources for the most important work and services according to policies and in response to changes
- *Failover and HA*: Detect service failures & dynamically switch them over to other available resources or sites
- *Service Aggregation*: Compose higher level services from other services; services supported by multiple instances in a data centers or across grid

Workload Management: Effective processing of various types of user work activities:

- *Distributed Batch*: Effective processing of non-interactive jobs across grid

- *Flow scheduling*: construction, organization, sequencing and staging of related jobs into operationalized flows according to dependencies and calendars.
- *Distributed Messaging*: Messaging workload processing across grid
- *Session Load Balancing*: Scheduling of interactive session-based applications onto servers across grid

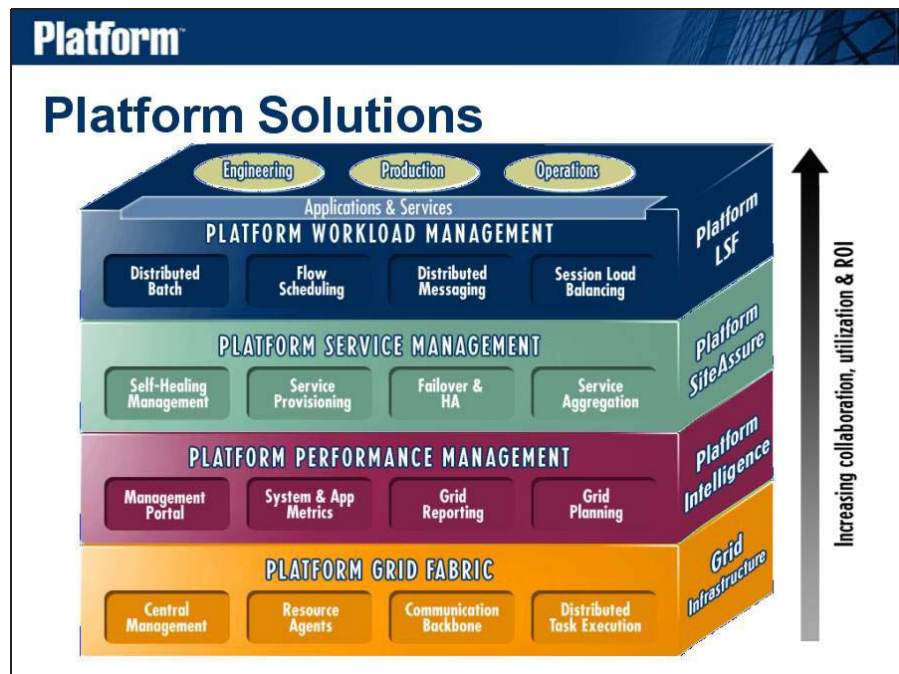


Figure 11



Figure 12



Figure 13

PLATFORM INTELLIGENCE

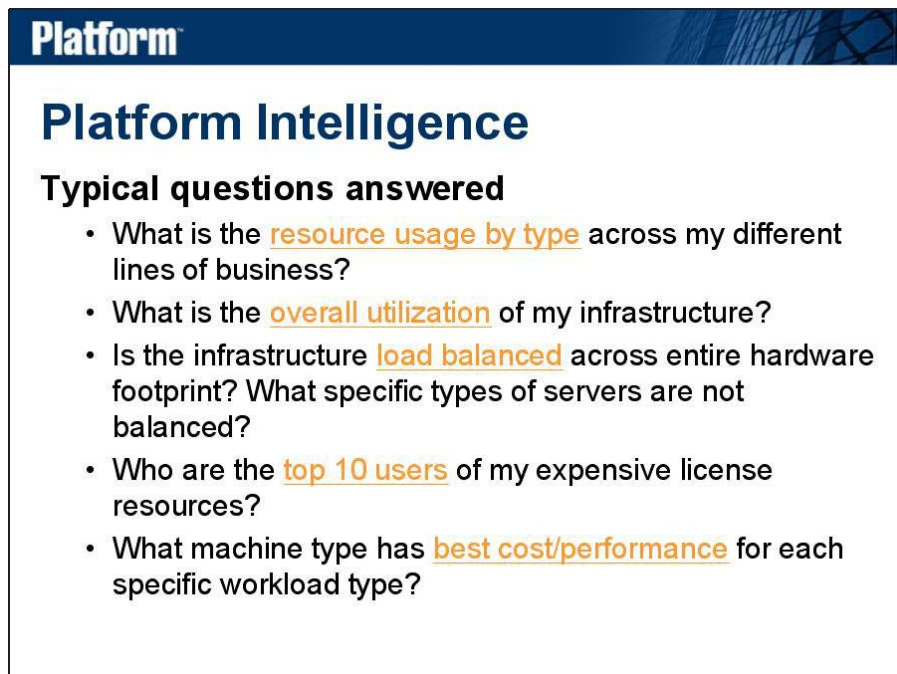
Here is an example of the types of business questions that can be answered with Platform Intelligence.

By geographic location or business unit recognize how systems and servers are deployed and used.

What are the utilization patterns across the whole infrastructure? Where is there head room? Can 20% of the services be reallocated on to existing servers?

What the cost of over utilized licenses? Are they performing productive work or more importantly, the productive work critical to the success of the company?

Where am I getting the best cost performance ratio? Is it on large SMPs? Desktops? Linux blades? What value would be placed on knowing these details?

A presentation slide titled 'Platform Intelligence' with a blue header. The slide lists 'Typical questions answered' with five bullet points. The text is white on a dark blue background for the header and title, and black on a white background for the list.

Platform

Platform Intelligence

Typical questions answered

- What is the **resource usage by type** across my different lines of business?
- What is the **overall utilization** of my infrastructure?
- Is the infrastructure **load balanced** across entire hardware footprint? What specific types of servers are not balanced?
- Who are the **top 10 users** of my expensive license resources?
- What machine type has **best cost/performance** for each specific workload type?

Figure 14

HOW IT WORKS

Platform Intelligence provides unique value by adding in the real world use of the systems and applications. By folding in areas such as locations, business units, projects and the people making use, it delivers real added value on top of the rich metrics concerning applications and systems.

Transforming, aggregating and building the data warehouse consistently on a daily or even hourly basis is a great challenge. With Platform Intelligence the steps of driving this data into the database and generating the resulting reports and OLAP cubes has been deeply enhanced and fine tuned. This allows Platform Intelligence to scale beyond 1500 hosts each tracking 20 metrics as is the case in our QA lab. Scalability is our business, so the boundary conditions are huge 10's of millions or row databases needing to be transformed for presentation.

The resulting OLAP analytics and reports are updated automatically using built in mechanisms. The data is refreshed constantly, so viewing and interacting with the data is always the most current and the most timely.

ROI = use + asset value + business objectives

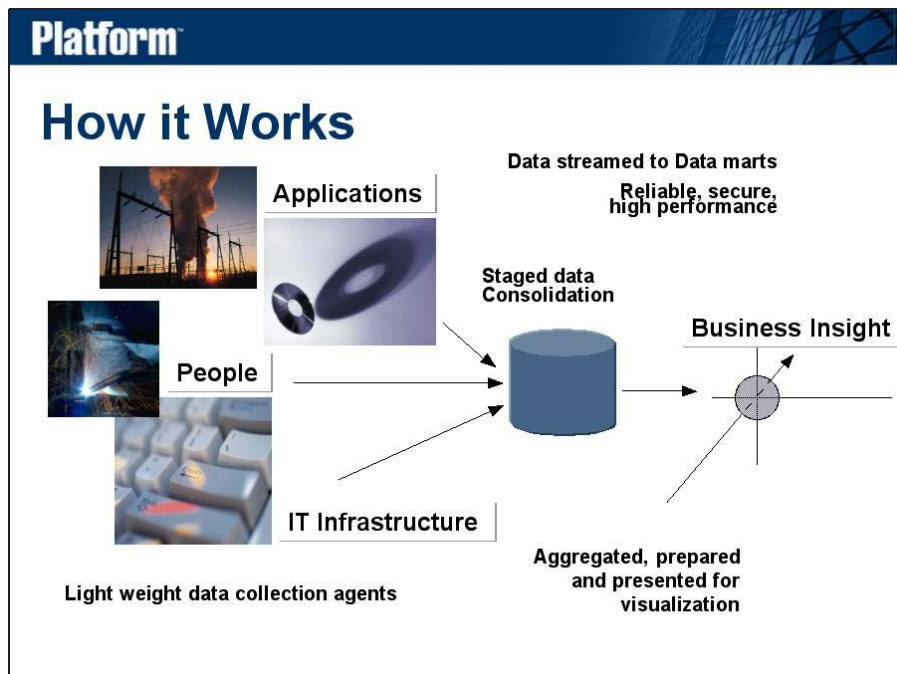


Figure 15

PLATFORM INTELLIGENCE

The layered architecture builds on the traditional view of business intelligence. At the bottom level, data collected, or extracted from correct sources. It is presented to the operational data store, with the metadata intact resulting in guaranteed consistency. By performing the required aggregations, the data is reduced in size, where possible without losing any insight. By performing these operations, the resulting visual tools will be smaller, load quicker and perform better.

From the operational data store (ODS) the OLAP cubes and reports are built. These multi-dimensional viewers provide a means to interact with the focused subject areas that are of most interest. Be it project, licenses, workload or IT performance are pre-organized to speed the access the the greatest value. By linking the cubes it is also possible to identify an area of focus and follow it through to another OLAP multidimensional view.

Users interact with Platform Intelligence through a management portal that is personalized for the individual user. Be operational administrator or executive, each can have a view tailored to their unique informational needs, that is also controlled by the user in order to continue to the modifications. This portal is based on Internet technology and supports secure access in a zero foot print fashion from any access point across the enterprise.



Figure 16

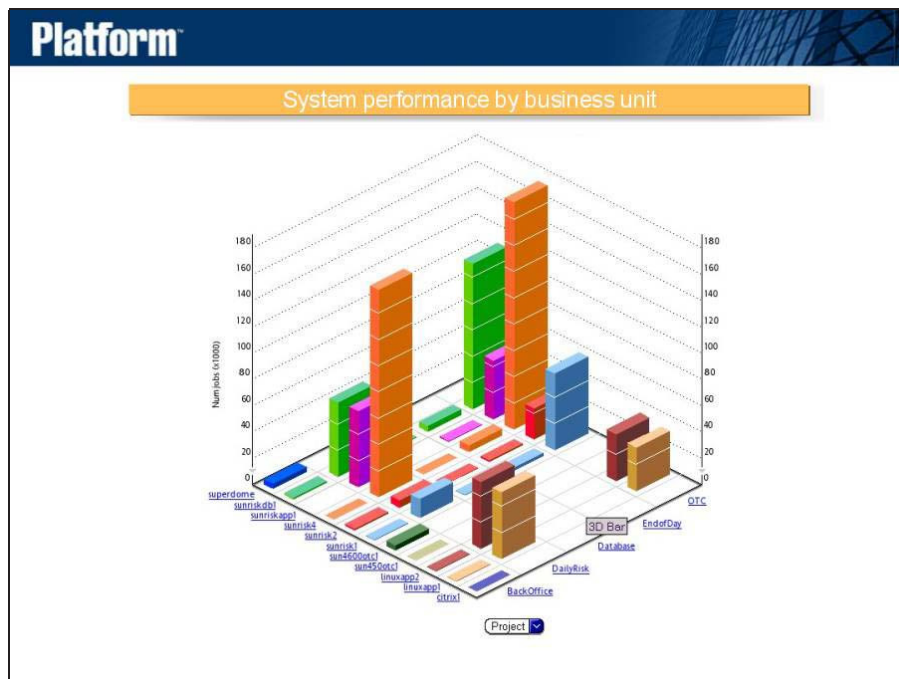


Figure 17

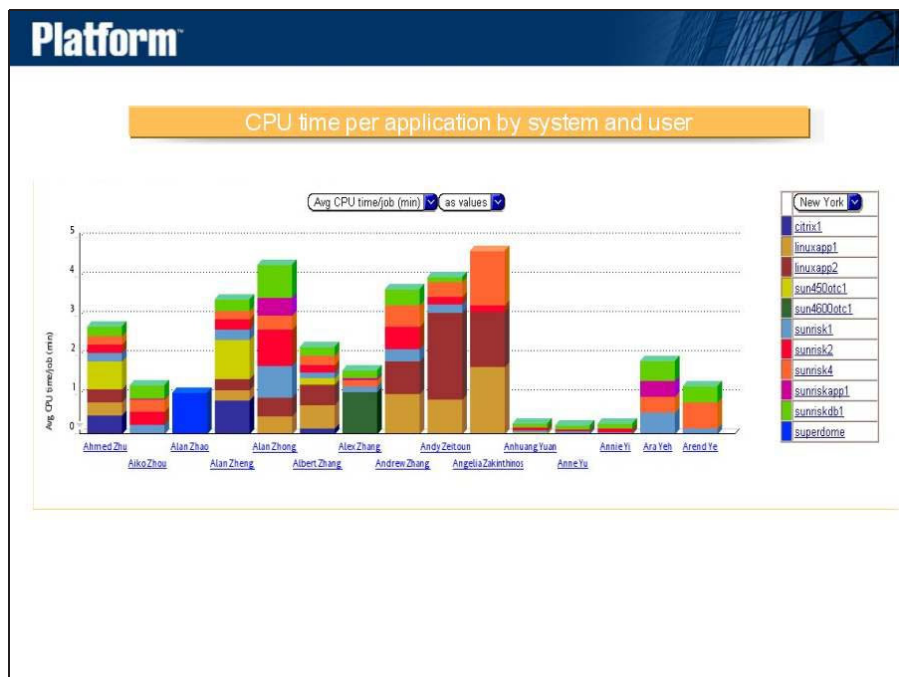


Figure 18

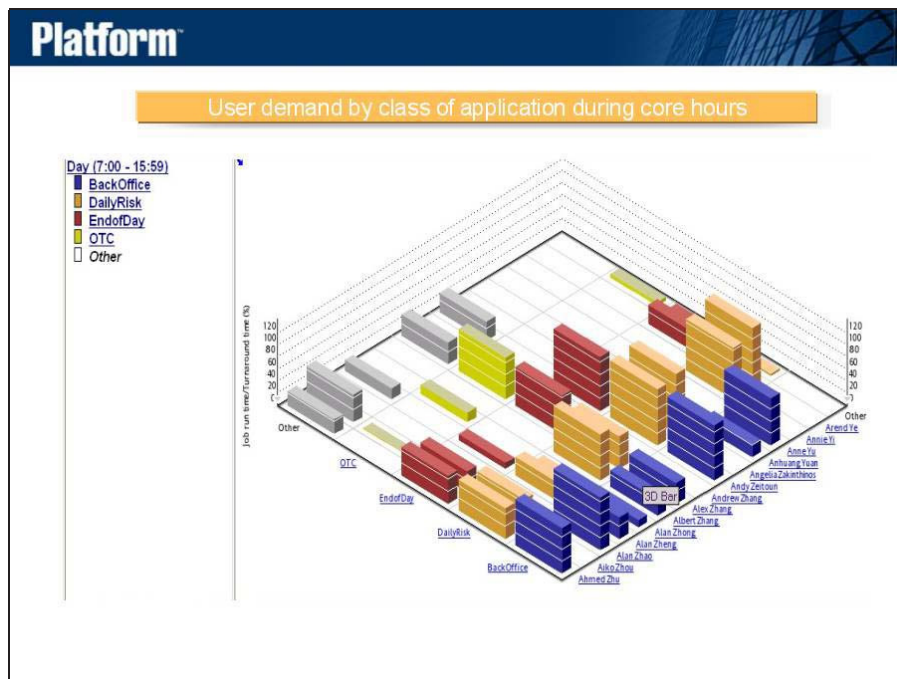


Figure 19

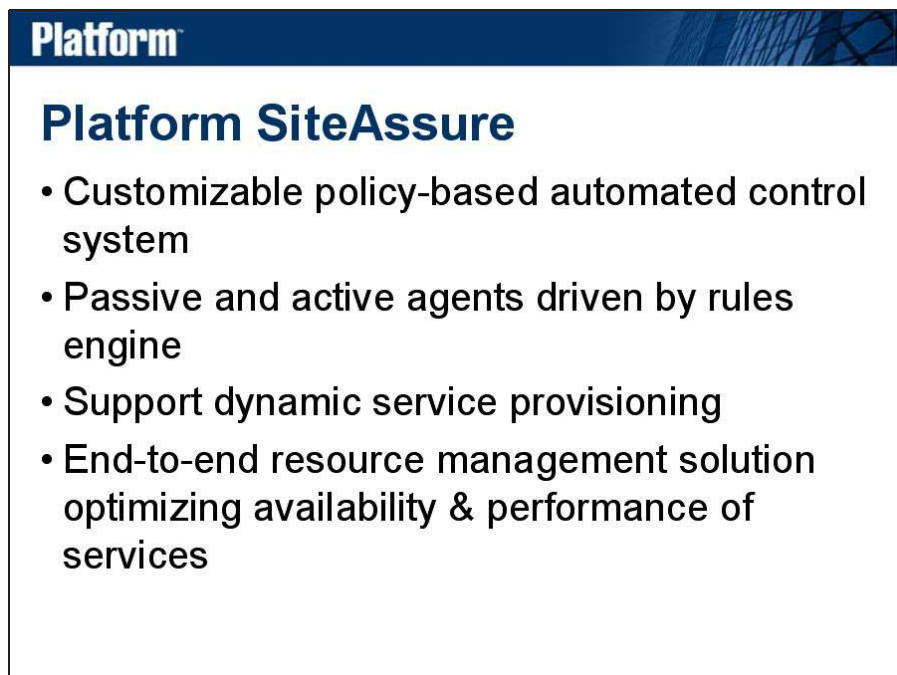


Figure 20

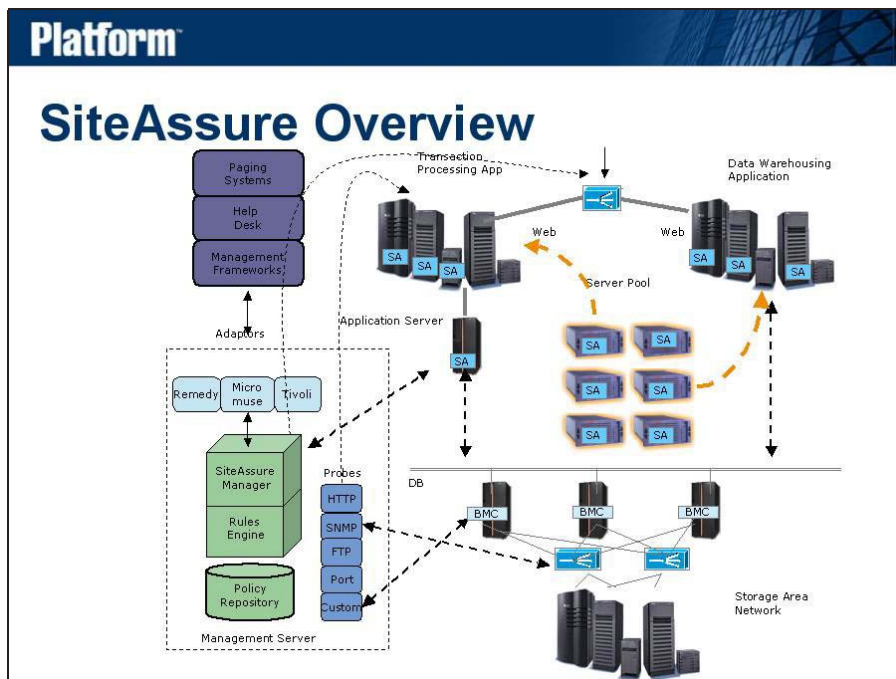


Figure 21

Platform LSF

"I want to make sure our users' work is processed reliably, timely, and easily - I want them to focus on their work, not IT"

- Work Load Management optimizes the productivity of the computing environment whereas Service Management manages computing supply
- Our value is distributed scheduling, reliable management, dynamic service allocation, heterogeneous resource matching, and application integration
- Enterprise workload requires more than a load balancer switch

Figure 22

OPEN, SCALABLE WORKLOAD MANAGER

Scalable, Grid-enabled to enable extensibility and customizability to meet growing needs of users. The mbatchd is split up into two processes: a manager process, and a scheduler process. The scheduler is further modularized into a number of plug-ins. The scheduler process loads a number of scheduler plug-ins corresponding to specific policies. Note that MultiCluster is handled as just another scheduler plug-in. The Manager and Scheduler communicates over a socket. The separation of the mbatchd into two processes enhances performance as it allows the scheduler to focus on scheduling while the manager can do the overall coordination, including handling client, LIM, and sbatchd interactions. Additionally, the scheduling data structure and algorithm has been restructured so that the scheduler is resource-centric; it maps resources to available jobs rather than vice-versa as was the case in 4.x. This means that as long as a resource is available, the scheduler will schedule the next waiting job. This improves the performance when there are large number of jobs in the system significantly as the scheduler does not need to go through a long list of jobs. The scheduler plug-in API is unique in that multiple plug-ins can co-exist at the same time, and can complement each other's policies. This is in contrast to SGE/PBS where the scheduler must be entirely replaced. A site can write additional scheduling policies and simply plug into LSF to complement the existing LSF policies. An LSF Web Service Broker is introduced to support SOAP/XML interface into LSF. This is consistent with our standards-based direction. The SOAP/XML interface means users can access LSF functionality (e.g., submit a job) programmatically in a platform-independent manner. The Web GUI is built using the LSF Web Service Broker. The GUI uses the .Net infrastructure.

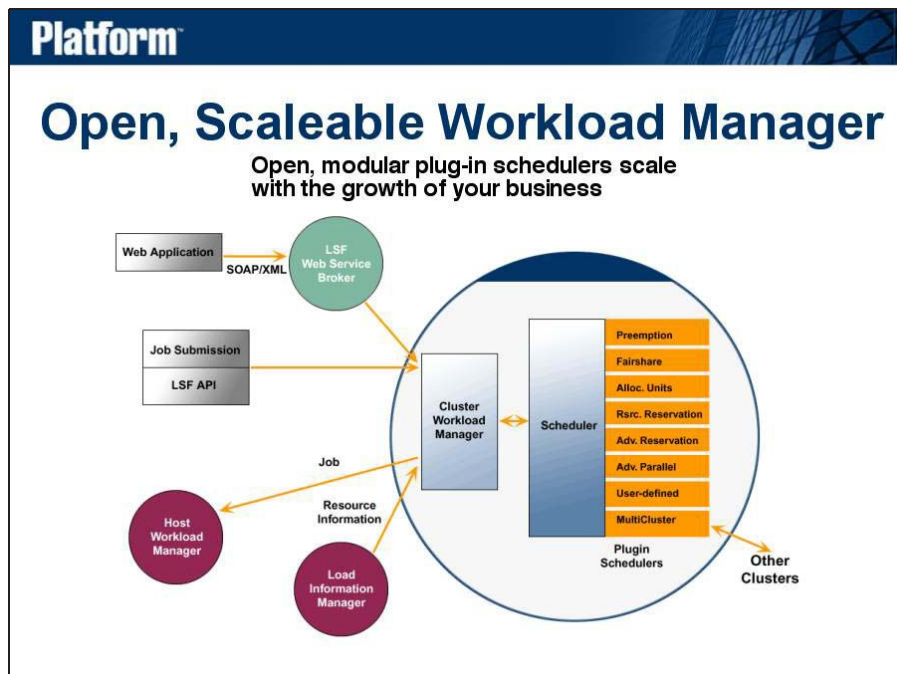


Figure 23

MULTICLUSTERS

With this Figure we want to highlight our complete end to end capability with Grid Computing to clusters to subclusters of desktops. NO ONE else has this complete solution for Life Sciences customers. With the growing computational requirements it is necessary to leverage all resources across an org.

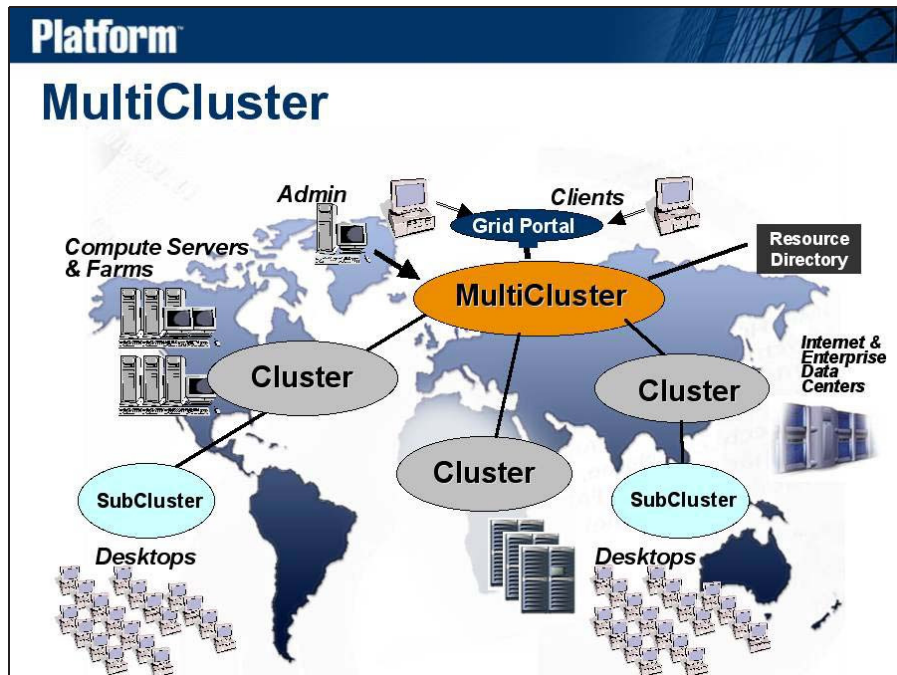


Figure 24

ABOUT OUR PRODUCT SUITE

In many cases we see companies consider the ‘workload management’ component of grid technology. While clearly the ‘bread and butter’ of grid computing, we consider three dimensions to a successful grid implementation: The tools to build, run, and manage a grid environment.

The development environment considers the tools required to build grid applications. In our product suite we have IDEs as well as GUI-based grid workflow processing design tools.

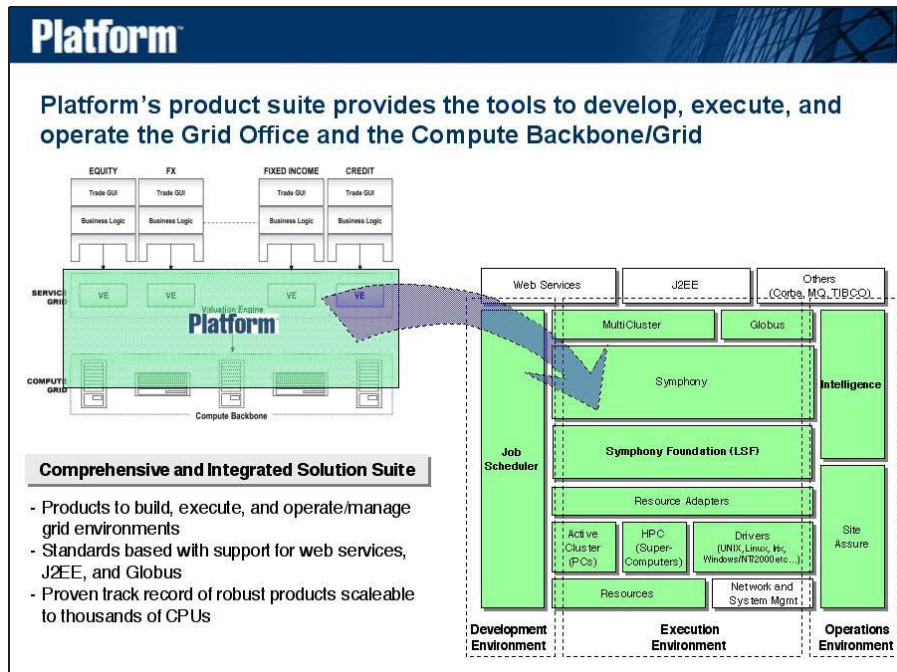


Figure 25

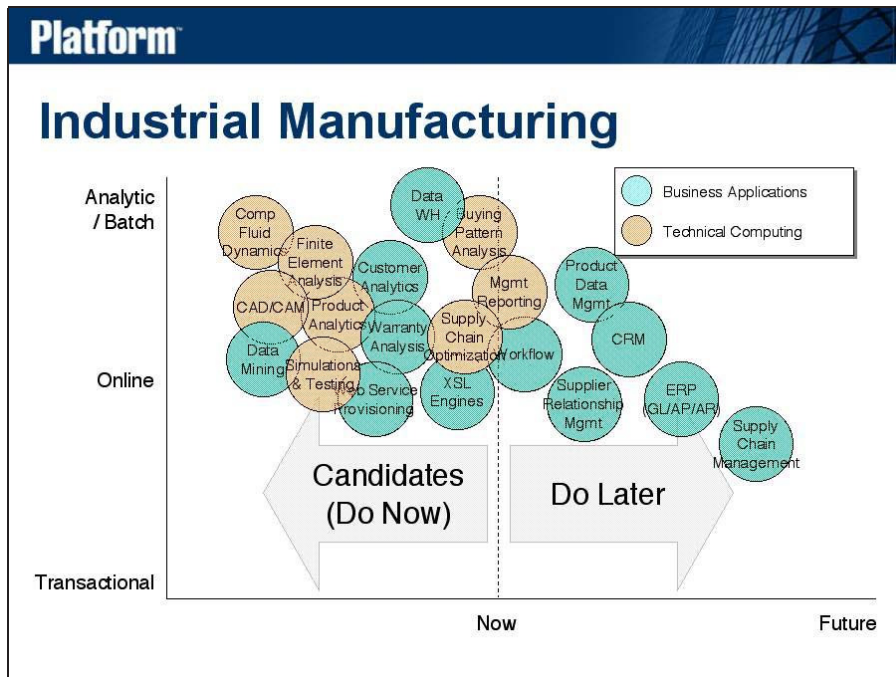


Figure 26



Figure 27

Platform

TACC

Texas Advanced Computing Center (TACC)

Overview

- Leader in inter-Grid collaboration
- TACC is building a University-wide Grid to connect clusters, workstations, visualization systems, and storage devices

Challenge

- Accelerate collaborative computational science and engineering on Grids
- Enable researchers to share HPC resources & execute codes across multiple HPC systems – Grid of Grids

Solution

- Platform LSF
- Platform MultiCluster
- Platform Globus

Results

- Web-based Grid portal simplifying use of many HPC systems
- Interoperability between local, state, national and global Grids

"By partnering with industry leaders like Platform, we hope to accelerate the collaborative nature of science on grids. Because TACC can leverage Platform's work with Platform Globus, Platform LSF, and Platform MultiCluster at the outset, we can focus our attention on the cutting edge of enabling science. Together, we can leverage existing research initiatives and develop new solutions rapidly. Much of the work we do today in universities is what industry will demand tomorrow."

Jay Boisseau,
Director, TACC

Figure 28

Platform

UTGrid Web Portal

Back Forward Stop Refresh Home Print Mail Add

Address: http://grid.tacc.utexas.edu/

TACC ATT Email TACC WASH Google UT Portal

University of Texas at Austin

Grid Computing Portal

Information

Available Systems

Grid Status

Job Status

File Manipulation

List Remote Files

List Portal Files

File Upload

Transfer to Remote

Transfer to Portal

3rd Party Transfer

Scientific Apps

Seismic Application

Demo Apps

PI Demo

Log In

Dept	System/Processors	Peak GFLOPs	Memory GBytes	Work. Disk GBytes	Name	Grid SW	Network	Status	Load	Jobs
CS	Linux PC	1.5	.1	52	alpha			↑		
CS	Linux PC	1.5	.1	52	solitude			↑		
TACC	Cray SV1 / 16	19	16	485	aurora			↑		71
TACC	Linux Cluster / 2	1	.5	13	braves			↑		
TACC	Linux PC	2	1	10	cool		Q	↑		
TACC	IBM Regatta-HPC / 64	333	128	532	longhorn			↑		4R-4Q
TACC	LSF Multi-Cluster / 22	37	14	173	luf			↑		6R-2Q-3Q
TACC	Linux Cluster / 4	2	1	13	padre			↑		
TACC	Cray/Dell Cluster / 4	19	8	8	q			↑		
TACC	Linux PC	2	1	10	sanantonio		Q	↑		
TACC	IBM IA-64 Cluster / 40	128	80	140	santarcia			↑		
TACC	Sun Workstation	2	1	2	tahoka			↑		
TACC	IBM IA-32 Cluster / 64	64	32	20	tejas		Q	↑		6R-4Q-2Q
TICAM	Alpha Cluster / 16	16	8	71	zaphod			↑		
	Total:	627	290	1581						

Click on column headers to sort.

Click the magnifying glass for more information about grid software status or network connectivity.

TACC

Internet cone

Figure 29

165

Platform™

Seismic LSF Leasing Job Submit

This method leases processors from padre if the job needs additional processors. Brazos has 2 processors and up to 4 can be borrowed from Padre.

LSF Leasing Job Submission	
Job Name:	Seismic_17-11-2002_19:50
Parameters:	fx: 0.0 dx: 6.0 fz: 0.0 dz: 6.0 xs: 0.0 zs: 0.0 tmax: 4.0 mt: 100 fpeak: 20.0 fmax: 40.0 verbose: 1 hsz: 0.0
Machine:	Brazos (UT Austin)
Number of Processors:	4
Directory for Results:	~/
Output File:	17-11-2002_19:50.out
Note: This may take a few minutes to run.	
Submit LSF Job	

TACC

Figure 30



Figure 31

Platform

ASCI (Accelerated Strategic Computing Initiative Grid) – Sandia National Laboratory

Overview

ASCI securely connect 3 geographically dispersed U.S. Department of Energy Labs

- Sandia National Labs (Albuquerque, NM)
- Lawrence Livermore National Lab (Livermore, CA)
- Los Alamos National Lab (Los Alamos, NM)

Challenge

Required Kerberos-compliant Grid solution

Solution

Platform Globus with Kerberos security

Results

Kerberos-compliant Partner Grid



"Platform Globus offers us numerous benefits, such as multi-platform commercial technical support, quality assurance and cost effectiveness. We made initial steps in hardening the Globus Toolkit for use with Kerberos, and are now collaborating with Platform to enhance the initial integration."

Steven Humphreys,
ASCI Grid Services Project

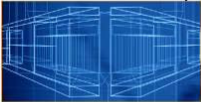


Figure 32

Platform

TRW

Overview

Leader in advanced technology for aerospace, information systems and automotive markets worldwide

Challenge

Increase license availability across multiple geographic & determine actual use per user

Solution

Platform Global License Broker
Platform Intelligence

Results

Transparently share software licenses across multiple geographic locations without manual intervention
Expect to reduce additional license purchases



"With this solution we can pre-empt idle licenses from interactive sessions, and track usage to determine actual use per user - all without human intervention."

Al Danial,
Manager, Engineering Applications,
TRW

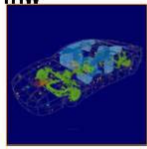


Figure 33

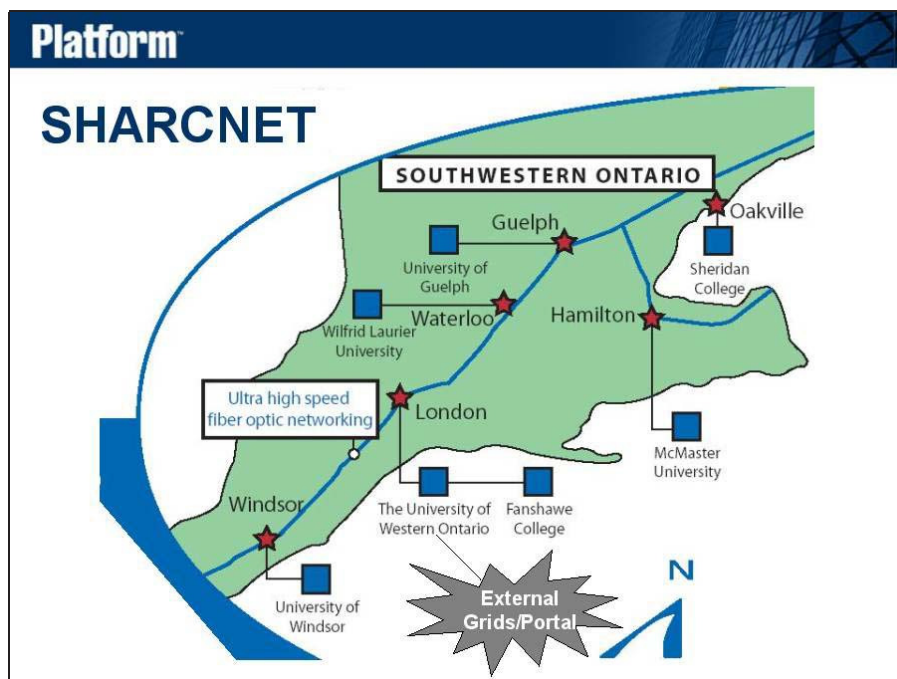


Figure 34

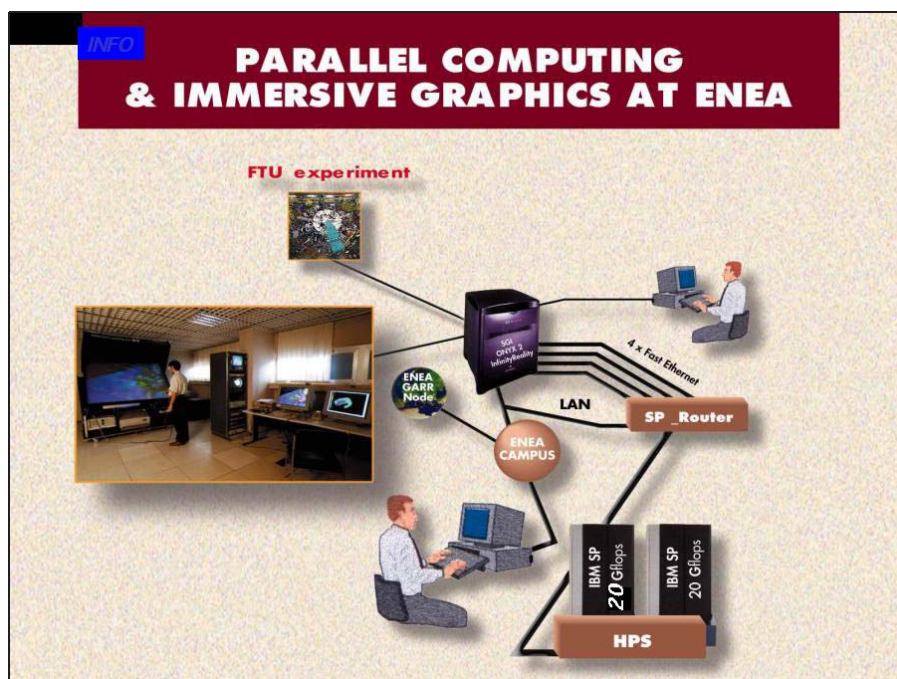


Figure 35

Grid at Pharmacia

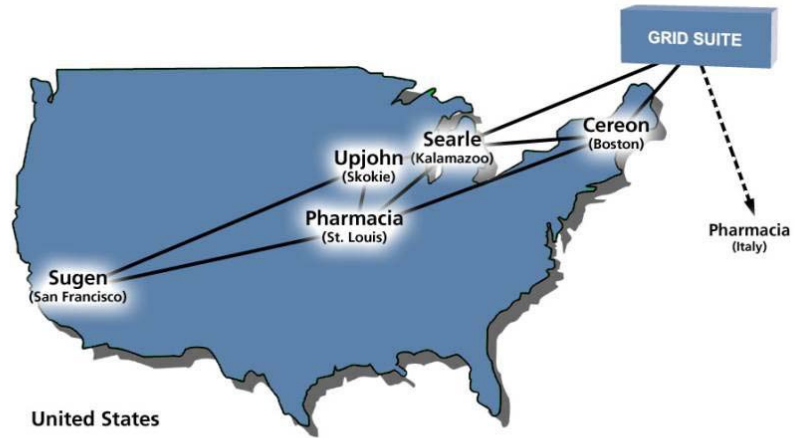


Figure 36

HOKKAIDO UNIVERSITY COMPUTING CENTER IN SAPPORO, JAPAN

The budding grid environment is based on the "e-science" concept of collaborating on and sharing research electronically, as recommended by Japan's Ministry of Education, Culture, Sports, Science and Technology.

Platform

Hokkaido University

- National public university
 - First of seven e-science centers
- Scientific disciplines
 - Environmental studies
 - Nanotechnology
 - Observational astronomy
 - Post-genomic bioinformatics
- Grid-enabled visualization
 - SGI Visual Area Networking
 - Platform LSF
 - Platform Globus

Platform

sgi
Japan

32-CPU SGI® Onyx® 300 visualization system

Figure 37

Grid-Enabled Visualization

- Scientists collaborate in real time
- Geography is not binding
 - Scientists don't need to relocate to the site of the compute, data or visualization services
- Required bandwidth is minimal
 - Thin clients move pixels, not data files
- Security is simplified
 - Authentication and authorization
 - Wire-level encryption (e.g. via SSL/TLS)

Figure 38

Grid at Texas Instruments

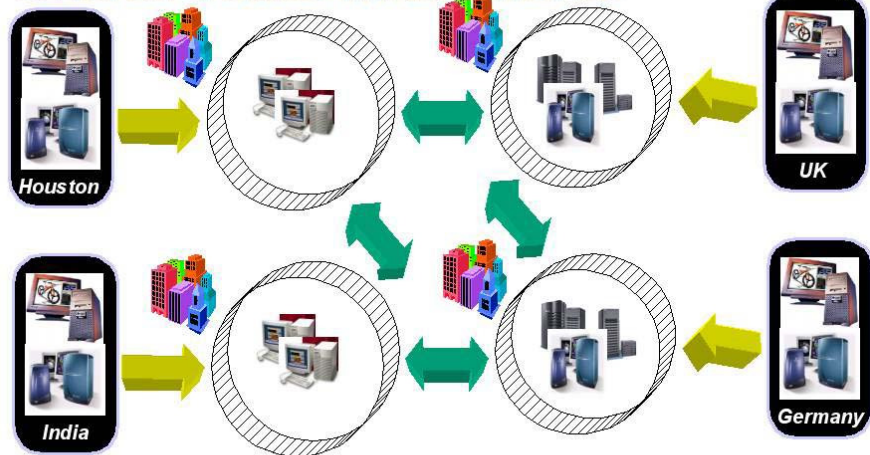


Figure 39

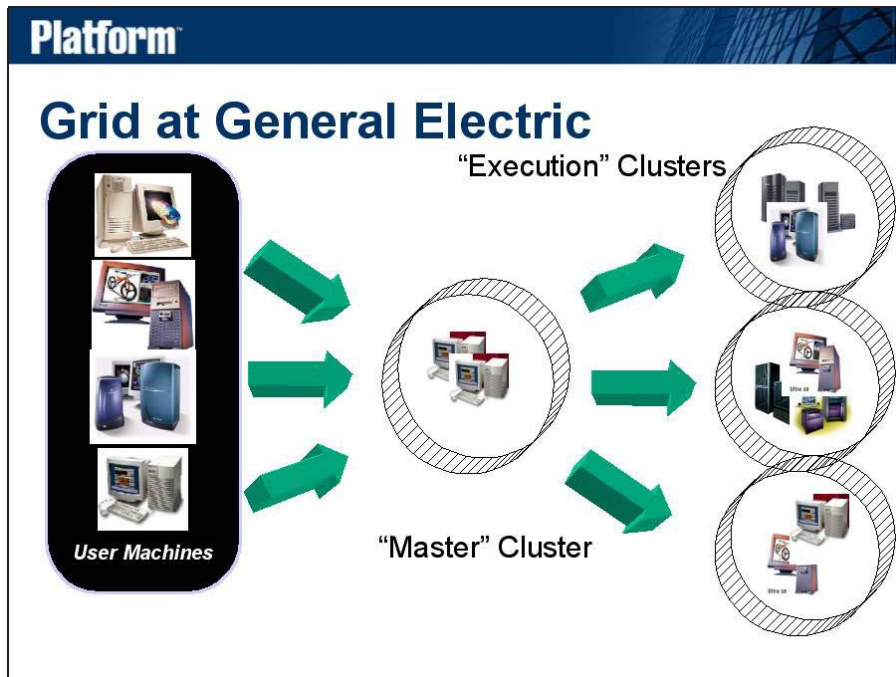


Figure 40

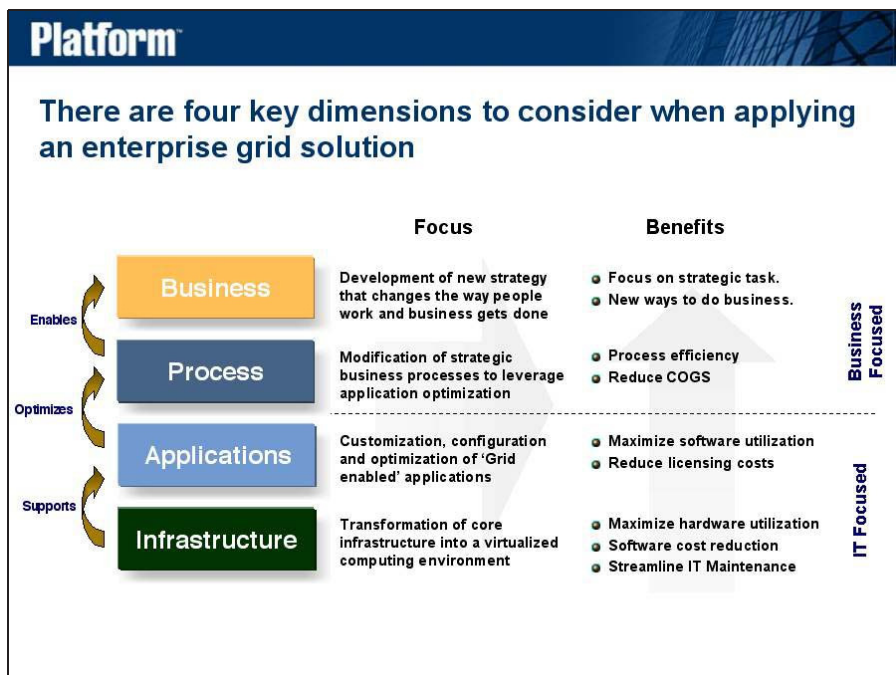


Figure 41

IPG POWER GRID OVERVIEW

Thomas Hinke
NASA Advanced Supercomputing (NAS) Division
NASA Ames Research Center
Moffett Field, CA

IPG POWER GRID OVERVIEW AND ACKNOWLEDGMENT

This presentation will provide a brief overview of the Information Power Grid.

I would like to acknowledge that many of the slides used in this presentation are based on a set of slides prepared by Tony Lisotta, for a grid tutorial that he recently presented at Global Grid Forum 7 in Tokyo.

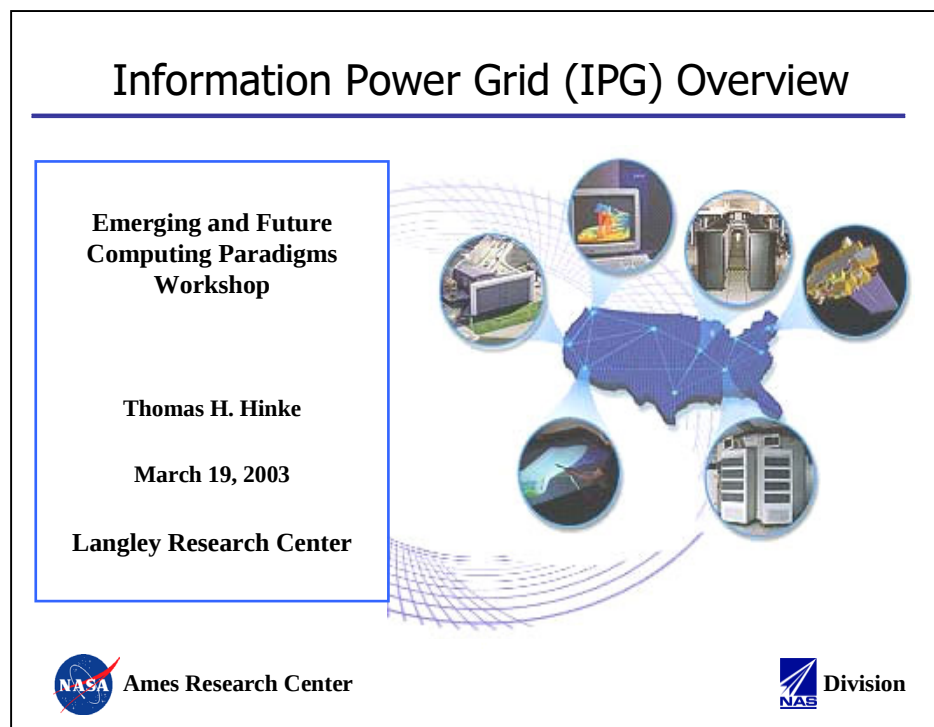


Figure 1

OUTLINE

This presentation will describe what is meant by grids and then cover the current state of the IPG. This will include an overview of the middleware that is key to the operation of the grid. The presentation will then describe some of the future directions that are planned for the IPG. Finally the presentation will conclude with a brief overview of the Global Grid Forum, which is a key activity that will contribute to the successful availability of grid components.

Outline

- What are Grids?
- Current State of Information Power Grid (IPG)
- Overview of IPG Middleware
- Future Directions
- Global Grid Forum

 Ames Research Center

 Division

Figure 2

WHAT DO GRIDS DO?

Grid software is middleware that sits on top of the network and the connected resources such as computers, storage and instruments. The grid software can provide an infrastructure on which to build collaborative environments that are large and distributed. They provide for security and provide the means to easily integrate distributed resources in a cost-effective manner.

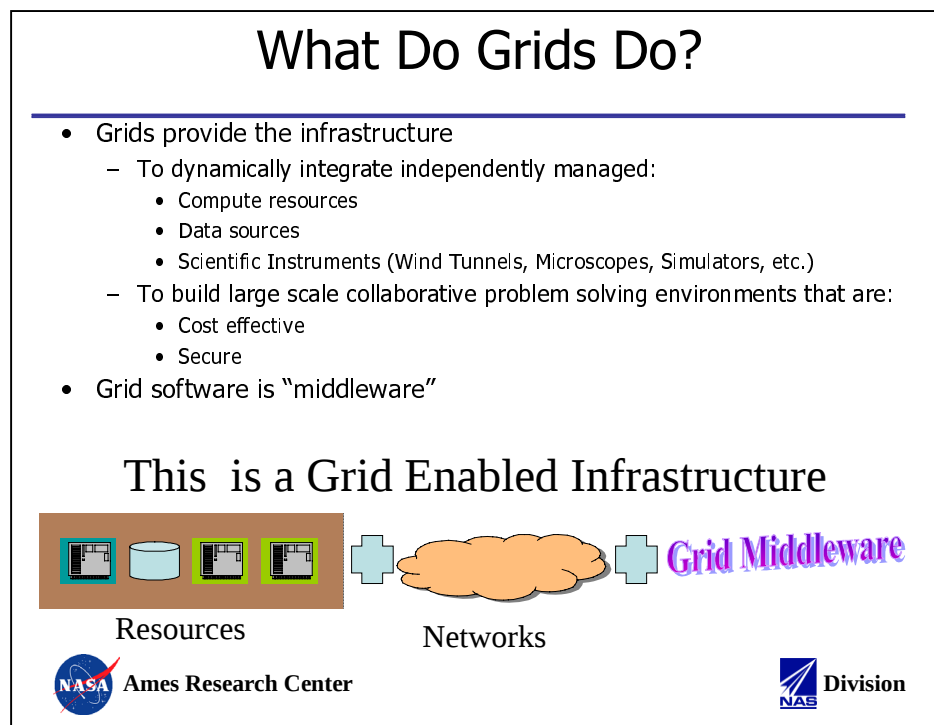


Figure 3

WHY USE GRIDS?

The goal of grids is to provide software that makes it easy for users to use distributed resources, such as distributed computers, storage or even instruments. The grid is actually a set of tools that permits these distributed resources to be easily accessed -- as if they were on the local system. These tools can also be used to develop distributed applications. They help the distributed application developer to focus on his applications, with the grid providing the software to handle the distributed access.

Why Grids?

For NASA and the general community today Grid middleware:

- Provides tools to access/use data sources (databases, instruments, ...)
- Provides tools to access computing (unique and generic)
- Is an enabler of large scale collaboration
 - Dynamically responding to needs is a key selling point of a grid.
 - Independent resources can be joined as appropriate to solve a problem.
- Provides tools for development of application-oriented frameworks
- Provides value added service to the NASA user base for utilizing resources on the grid in new and more efficient ways

 **Ames Research Center** **Division**

Figure 4

WHAT CHARACTERISTICS ARE NORMALLY FOUND IN A GRID

- Security is a fundamental aspect of a grid, with most grids basing their security on public key technology, which it used to protect at least the authentication information as it flows between the various sites on the grid. The IPG uses the Grid Security Infrastructure (GSI), based on the Globus toolkit, for its security.
- Using GSI, grids can support single sign-on, which means that after a user signs on one grid resource for a session, he is able to use other grid resources, on which he has an account, without any further identification or authentication required.
- Grids also provide a grid information service (GIS), that provides a single mechanism by which users can discover grid resources and associated information about the resource.
- Grids are designed to be scalable to a large number of resources.
- Finally, grids are designed to provide access to resources that may be under the control of different administrative groups. They are not designed to have centralized control.

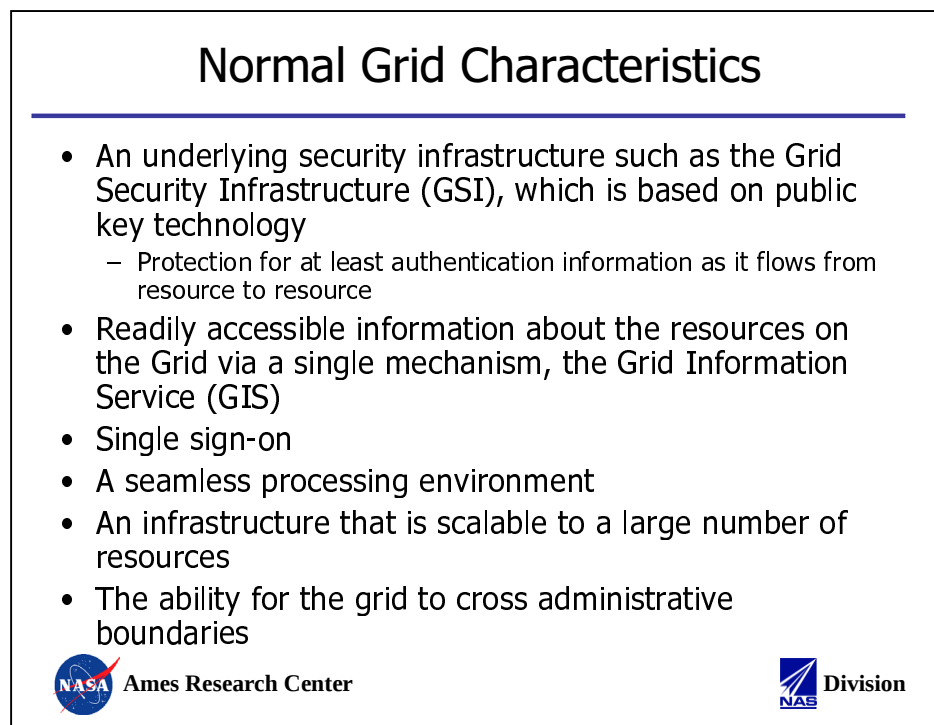


Figure 5

DISTRIBUTED SYSTEMS BEFORE THE GRID

Before the development of the grid, people still developed distributed systems. Under these pre-grid distributed systems, a user was responsible for dealing with all of the complexities of the distributed environment.

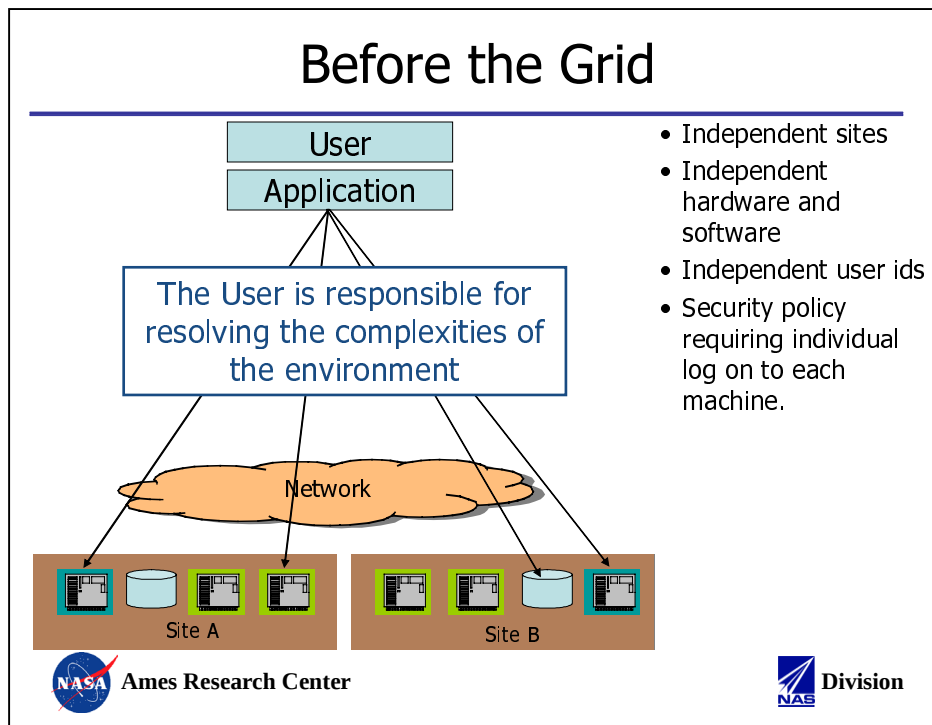


Figure 6

DISTRIBUTED SYSTEMS USING TODAY'S GRID

The grid provides the middleware that ties distributed resources into a seamless environment. Using the grid, a user can make a request to the grid Information Service for information about the location and characteristics of grid resources such as processing and storage resource or instruments. With this information, the user can then launch an application that accesses the desired distributed resources through the grid middleware.

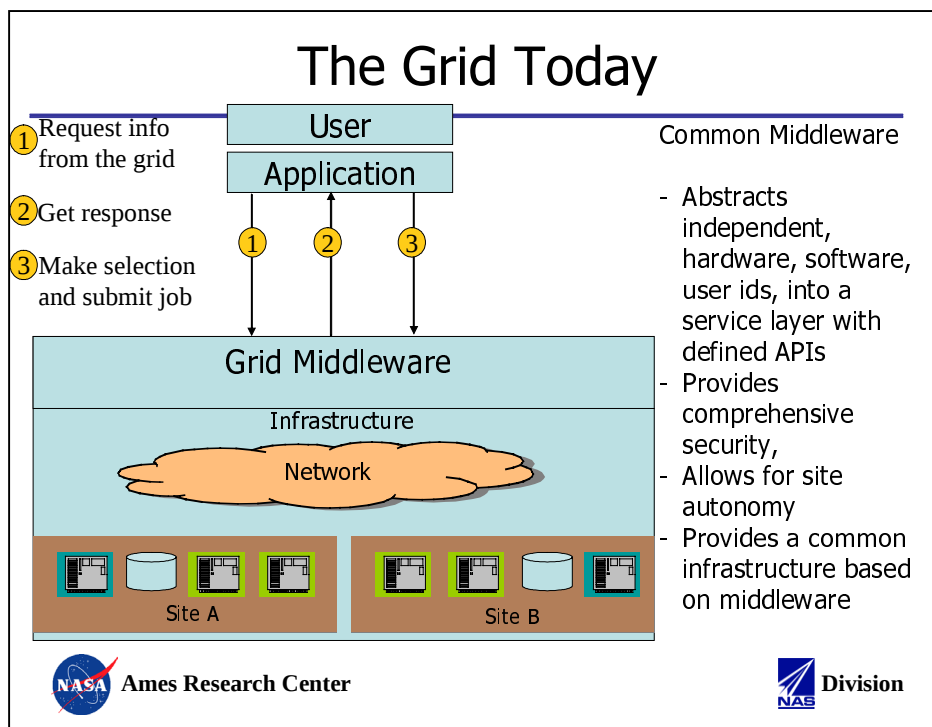


Figure 7

DISTRIBUTED SYSTEMS USING TODAY'S GRID

The key to the grid is that the underlying grid resources are abstracted into application programmer interfaces that simplify the development of distributed applications. While this is a significant step forward, this layer does not have much intelligence, which will define the next stage of grid development.

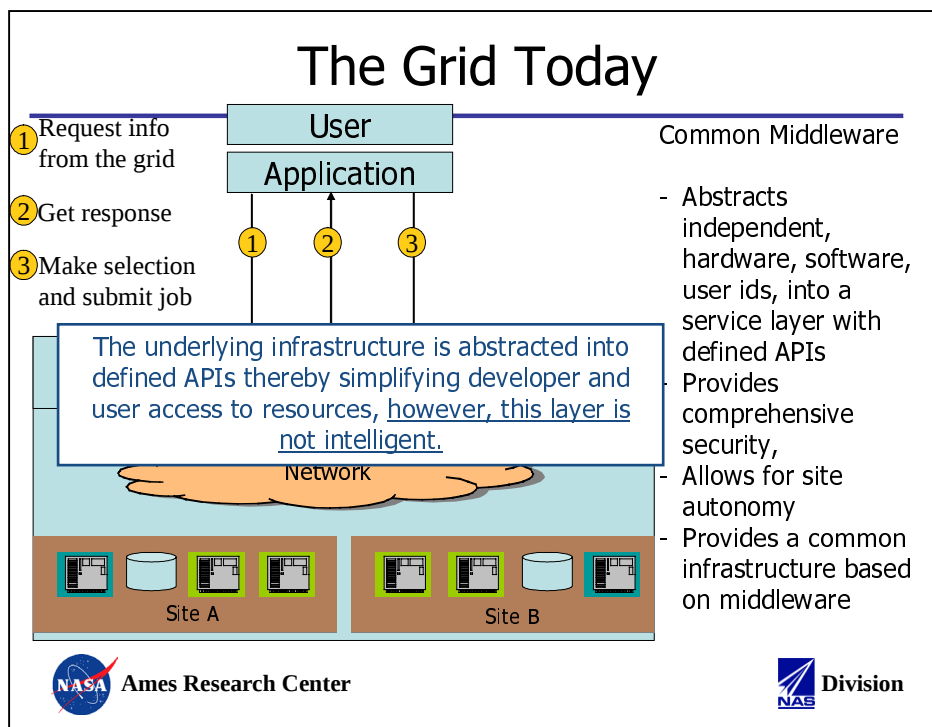


Figure 8

THE NEAR FUTURE GRID WILL HAVE INTELLIGENCE

The grid for the near future will have intelligent, customizable middleware that will sit between the current grid middleware and the application. This intelligent layer will perform brokering (the automatic selection of resources) and will provide information tailored to the specific needs of the user or application.

Under the current grid, a user must have an account on each resource that is used, thus preserving local autonomy. Under the near future grid, if a local system agrees, the grid will then take responsibility for granting grid user's access to these resources, where the user has not pre-established an account.

Another key capability that will soon be available is the ability to field grid-enabled web services, that provide a standard API that can be accessed from applications, application-specific portals or command-line functions.

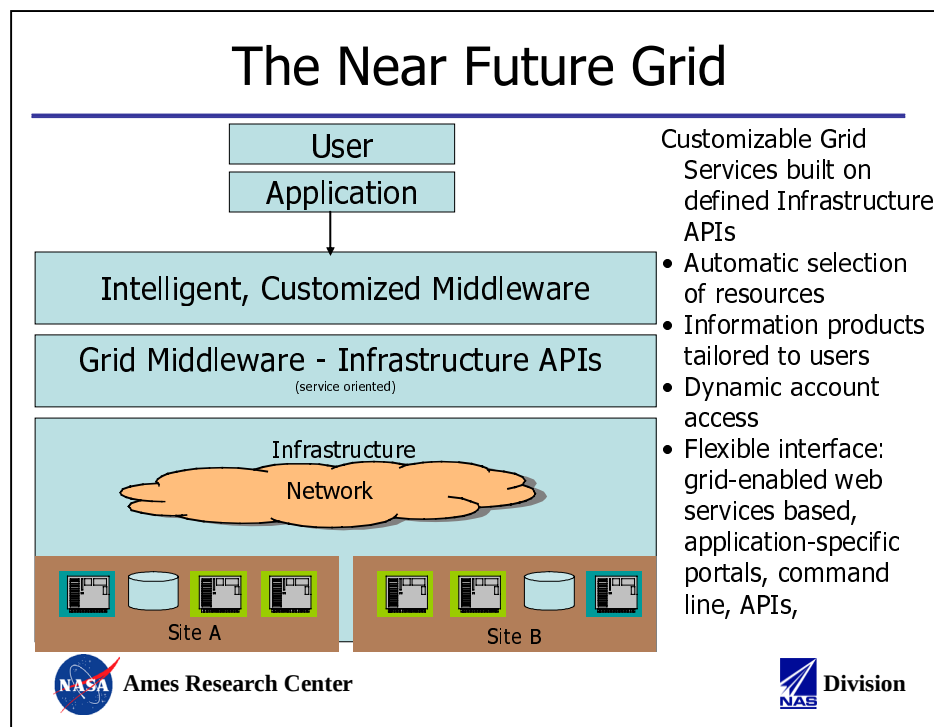


Figure 9

THE NEAR FUTURE GRID WILL HAVE INTELLIGENCE

With this more intelligent grid, the users and application developers will be able to focus more on the science and engineering applications and not on the distributed systems management aspects of their systems.

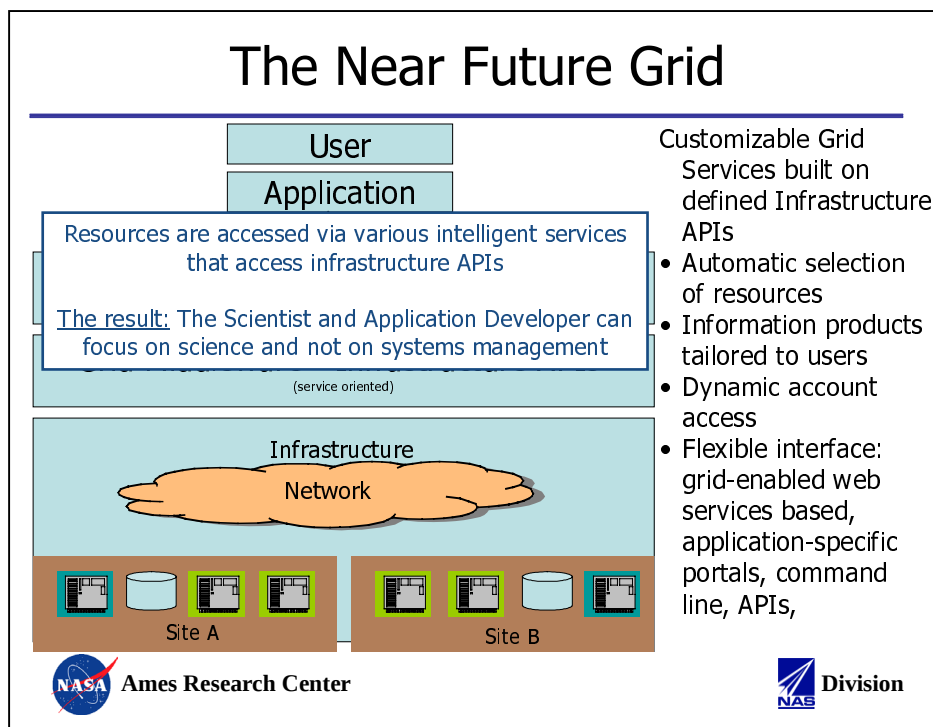


Figure 10

HOW THE USER AND APPLICATION DEVELOPERS SEE A GRID

A grid is really just a set of tools that can be accessed through application programmer interfaces or command line functions. These tools will be augmented with services that will be structured as grid-enabled web services, which are re-usable such that one or more of these can be combined to make a more complex services.

Once a user has authenticated to the grid, he can use any of the various services that are shown on the slide as if these were part of his local machine. He does not have to re-authenticate to use any of these, with the grid handling the requirement to pass identification and authentication information among the resources that are used.

How the User and Application Developers See a Grid

- A set of grid functions that are available as
 - Application programmer interfaces (APIs)
 - Command-line functions
 - Grid-enabled web services
- After authentication, grid functions can be used to
 - Spawn jobs on different processors with a single command
 - Access data on remote systems
 - Move data from one processor to another
 - Support the communication between programs executing on different processors
 - Discover the properties of computational resources available on the grid using the grid information service
 - Use a broker to select the best place for a job to run and then negotiate the reservation and execution (coming soon).

 **Ames Research Center** **Division**

Figure 11

OUTLINE

In the next section we will look at the current state of the IPG.

Outline

- What are Grids?
- Current State of Information Power Grid (IPG)
- Overview of IPG Middleware
- Future Directions
- Global Grid Forum

 Ames Research Center

 Division

Figure 12

IPG LOCATIONS

The IPG currently has resources located at the five NASA Centers shown on the map.

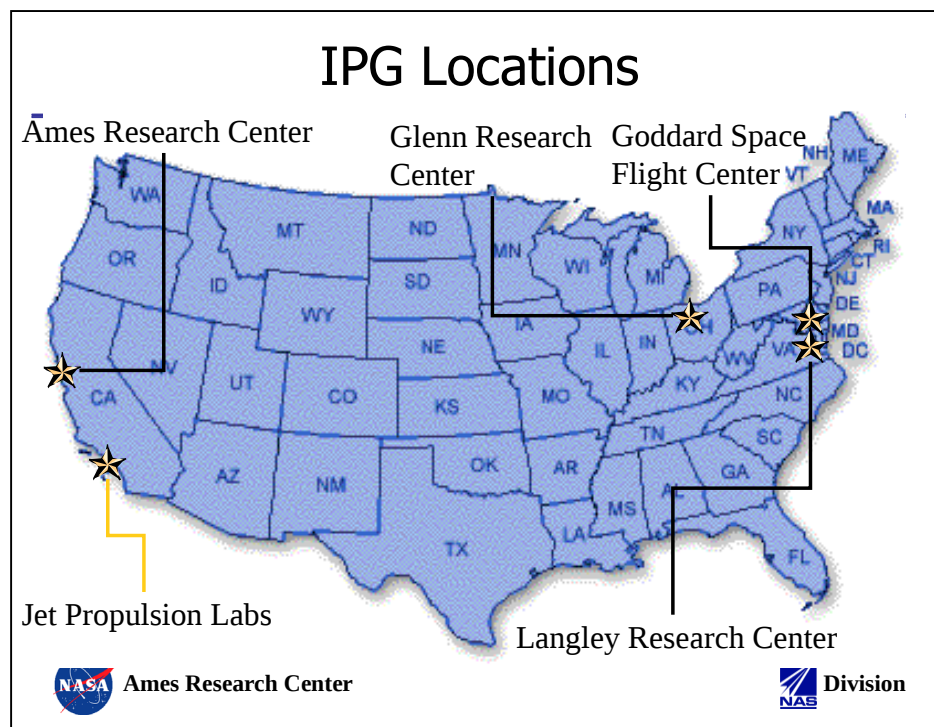


Figure 13

IPG RESOURCES

The IPG currently has the computational resources shown.

IPG Resources

- Server Nodes
 - 1024 CPU, single system image SGI, **Ames**
 - 512 CPU SGI O2K, **Ames**
 - 128 CPU Linux Cluster, **Glenn**
 - 124 CPU SGI O2K, **Ames**
 - 64 CPU SGI O2K, **Ames**
 - 24 CPU SGI O2K, **Glenn**
 - 16 CPU SGI O2K, **Langley**
 - 16 CPU SGI O2K, **Ames**
 - 8 CPU SGI, O3K, **Langley**
 - 4 CPU SGI O2K, **Langley**
- Client Nodes
 - 16 CPU SGI O300, JPL
 - 8 CPU SGI O300, Goddard
- Wide area network interconnects of at least 100 Mbit/s

 Ames Research Center Division

Figure 14

OUTLINE

The next section will delve more deeply into the nature of the IPG middleware.

Outline

- What are Grids?
- Current State of Information Power Grid (IPG)
- Overview of IPG Middleware
- Future Directions
- Global Grid Forum

 NASA Ames Research Center

 NAE Division

Figure 15

IPG IS BUILD ON GLOBUS TOOLKIT 2

The IPG, as are most of the grids in the world, is built on Globus Toolkit 2 (GT2). The Grid Security Infrastructure (GSI) is based on X509 certificates, secure socket layer (SSL) and Transfer Layer Security (TLS). This supports a GSI-enabled Secure Shell (SSH) and GridFTP (a high performance GSI version of FTP).

The Grid Information Services is based on LDAP (lightweight Directory Access Protocol) which supports the Monitoring and Discovery Service (MDS), which provides a directory of grid resources and attributes.

Finally, the remote execution of jobs is supported by the Globus Resource Allocation Manager (GRAM), which provides an interface to various batch schedulers (e.g., PBS and LSF), as well as systems that permits users to directly execute jobs via fork. It permits the launching of remote jobs.

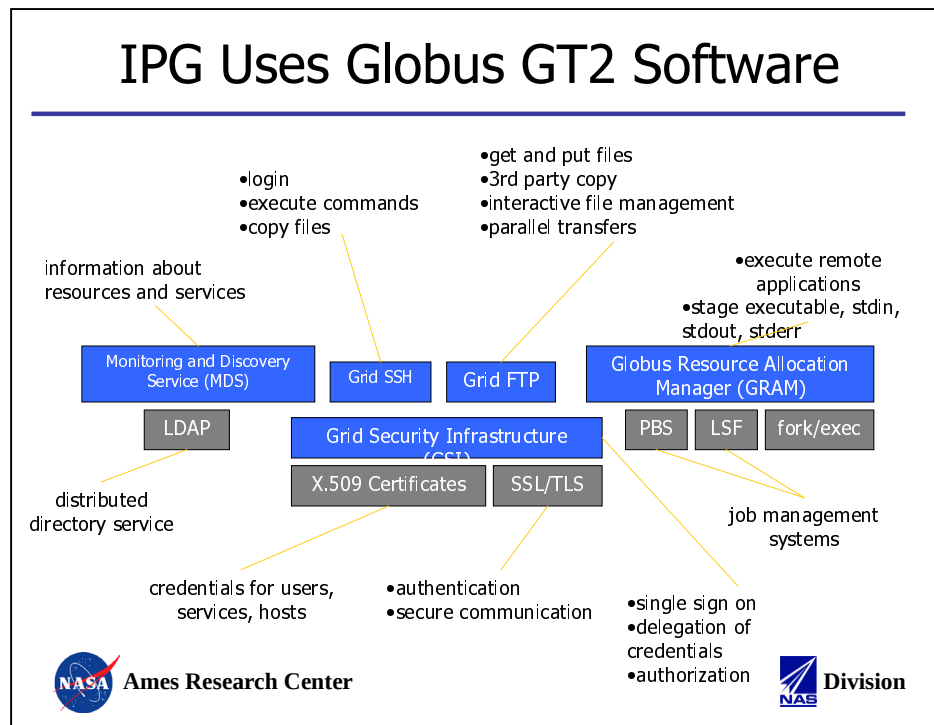


Figure 16

IPG/GLOBUS DEPLOYMENT ARCHITECTURE

To support the grid information service of a deployed grid, a Grid Resource Information Service (GRIS) captures local information from each resource and forwards this to a Grid Index Information Service (GIIS), that provides a single source for information about a particular grid.

Users, applications or web portals can use Globus client services to access any of the grid tools and services.

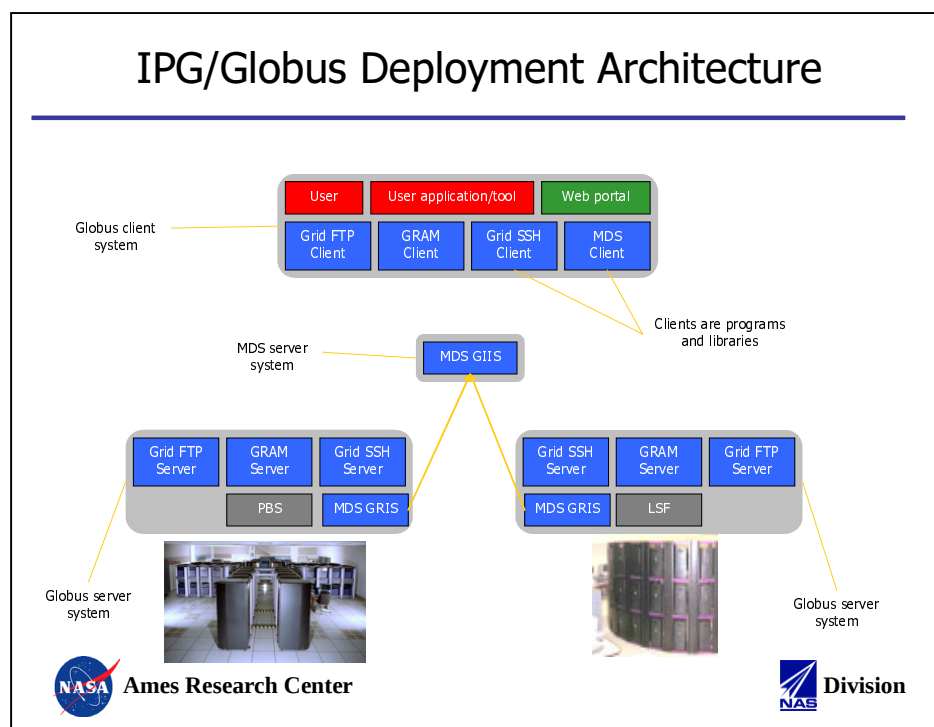


Figure 17

ADDITIONAL SERVICE UNDER DEVELOPMENT BY THE IPG PROJECT

To provide the added intelligence needed to facilitate the development of grid applications and the use of the grid by users, the IPG project is developing a Job Manager to manage the reliable execution of a job on the grid. The Job Manager will stage the necessary files needed by the application, monitor the progression of the work and then post-stage the results, cleaning up any files that may remain from the execution.

The Job Manager is supported by the Resource Broker that provides the user with suggestions about where to run his application, based on supplied information about the application.

Additional IPG Services

- Job Manager
 - Reliably execute a job
 - Set of files to pre-stage
 - Executable to run
 - Including directory, environment variables
 - Set of files to post-stage
- Resource Broker
 - Provide suggestions on where to run a job
 - Input
 - Which hosts and operating systems are acceptable
 - How to create a Job Manager Job for a selected host
 - Selection made using host and OS constraints and host load
 - Interactive system: # free CPUs
 - Batch system: Amount of work in queue / # CPUs
 - Output
 - Ordered list of Job Manager Jobs (suggested systems)

 **Ames Research Center**

 **Division**

Figure 18

Applications will be able to consult the broker for suggestions as to the best grid resources to use, given the current workload on each of these resources. This information will then be used to run the application on the suggested resources, using the job manager to stage necessary files and monitor the progress of the work and then post stage any files at the end of the work.



OUTLINE

Next we will briefly look at future directions.

Outline

- What are Grids?
- Current State of Information Power Grid (IPG)
- Overview of IPG Middleware
- Future Directions
- Global Grid Forum

 Ames Research Center

 Division

Figure 20

OPEN GRID SERVICES ARCHITECTURE (OGSA)

The Open Grid Services Architecture is the grid community's adoption of the web services work (which other than the name has little to do with the web) as a way of delivering services. Grid-enabled web services provide a standard Web Services Description Language (WSDL) description of the service and a specified protocol, which for now is SOAP, for accessing these services. Grid-enabled web services provide a self-describing way to offer services that can be included as components of other grid-enabled web service.

Standards are under development by the Global Grid Forum to specify the interfaces and the nature of the service-management capabilities (creation, destruction, lifetime) that are to be associated with each service.

One of the key contributions that grid-enabled web services offer over web services is that they will be built to use grid security, such as the Grid Security Infrastructure.

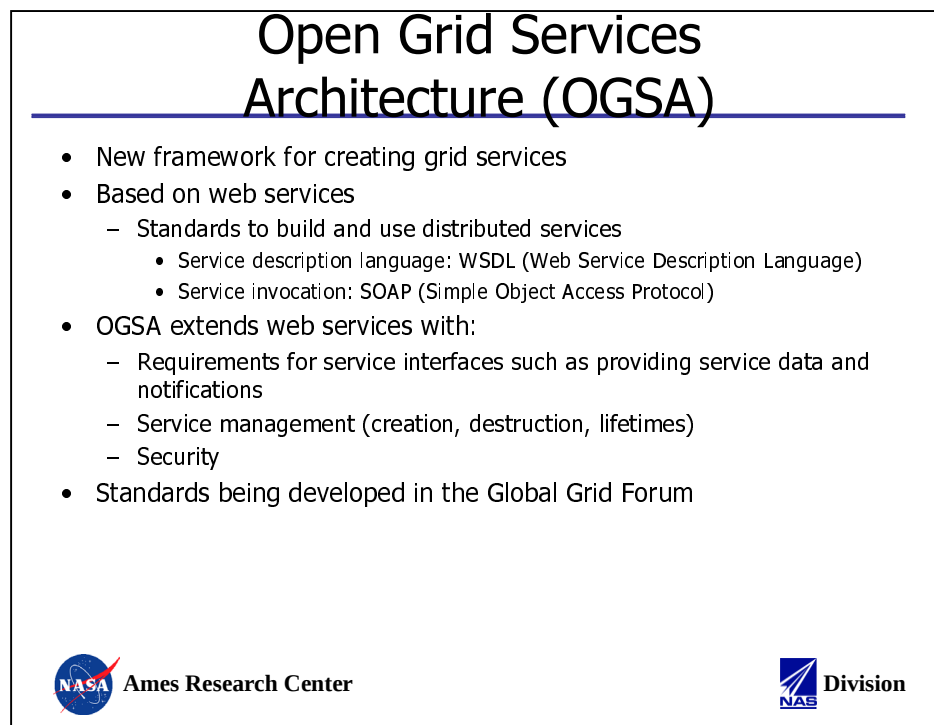


Figure 21

GLOBUS TOOLKIT VERSION 3 (GT3)

A key first application of OGSA will be the next version of the Globus Toolkit, which is called Globus Toolkit Version 3 (GT3). The various grid services offered by the Globus Toolkit will be offered as grid-enabled web services.

GT3 and OGSA will revolutionize how services are offered on the grid, since it will make it easy to include existing services in more complex, application-specific services.

The IPG will transition to GT3 as soon as it is stable and in a way that minimizes any impact to existing users.

Globus Toolkit Version 3 (GT3)

- Large change from GT2 to GT3
 - New implementation
 - Java-based instead of C-based
 - GT3 based on OGSA
- GT3 will provide equivalent services to GT2
- Alpha version of GT3 currently available
- GT3 and OGSA will revolutionize
 - how services are provided on the grid and
 - how grid applications are developed
- IPG will transition to GT3 soon as it is proven stable, while minimizing the effect on existing IPG users.
- Transition should have minimal impact on IPG users
 - Globus will maintain many of the existing programs
- IPG Services will follow OGSA

 **Ames Research Center** **Division**

Figure 22

FOCUS ON IPG HANDLING OF DATA

As the IPG completes its work on the resource management and utilization phase of the grid services, it will focus on the data handling aspects of the grid. This is a critical function for NASA because of the large volume of distributed data that is found in the various NASA archives, such as those associated with Earth science.

This new focus will look at providing access to NASA archives, using such existing grid-enabled systems as the Storage Resource Broker, developed at the San Diego Supercomputing Center. Of particular interest will be providing access to data stored on both tertiary storage (mass storage systems) and data stored on disk-resident data pools.

This effort will build on the considerable amount of work that has been performed on data grids by the international grid community.

Focus on IPG Handling of Data

- Goal: Intelligently manage data in a grid
- NASA data is inherently distributed e.g., various Earth science archives, including the one at LaRC
- Important focus of IPG
- Access to files
 - Initial use of grid-enabled Storage Resource Broker
 - Data staging and replica management building on grid community research
 - Need grid support for file metadata
- NASA data can be on
 - Disk-resident data pools
 - Tertiary storage data archives
- Will build on considerable data grid work from the international grid community

 Ames Research Center Division

Figure 23

OUTLINE

The last section will focus on the Global Grid Forum.

Outline

- What are Grids?
- Current State of Information Power Grid (IPG)
- Overview of IPG Middleware
- Future Directions
- Global Grid Forum

 NASA Ames Research Center

 NAE Division

Figure 24

GLOBAL GRID FORUM BACKGROUND

The Global Grid Forum is an international group that mirrors for grids what the Internet Engineering Task Force (IETF) has done for the network through its standards work. It was formed in 2001 as a combination of similar grid work in the North America and Europe and now encompasses the Asia/Pacific grid work as well. It meets three times a year in different parts of the world.

Global Grid Forum Background

- Began in 2001 as merger of previous regional grid forums.
- Now includes grid technical communities in North America, Europe and Asia Pacific
- Meets three times per year, alternating between North America and Europe and Asia/Pacific
- Modeled after IETF (Internet Engineering Task Force), which sets Internet standards.
- GGF7 was just held in Tokyo, Japan with over 700 attendees
- GGF8 will be held in Seattle, WA in June 25-27, 2003



Ames Research Center



Division

Figure 25

GLOBAL GRID FORUM PURPOSE AND ORGANIZATION

The main purpose of the Global Grid Forum is to provide an international grid organization that can support the fair and representative development, review, approval and release of both best practices and standards for the grid.

It is organized into two types of groups. The Working Groups are of limited duration and are focused on the goal of producing some specific best practice document or standard. Currently there are 24 Working Groups.

The Research Groups are organized to address grid issues that are not yet ready for a best practice document or a standard. Currently there are 20 research groups.

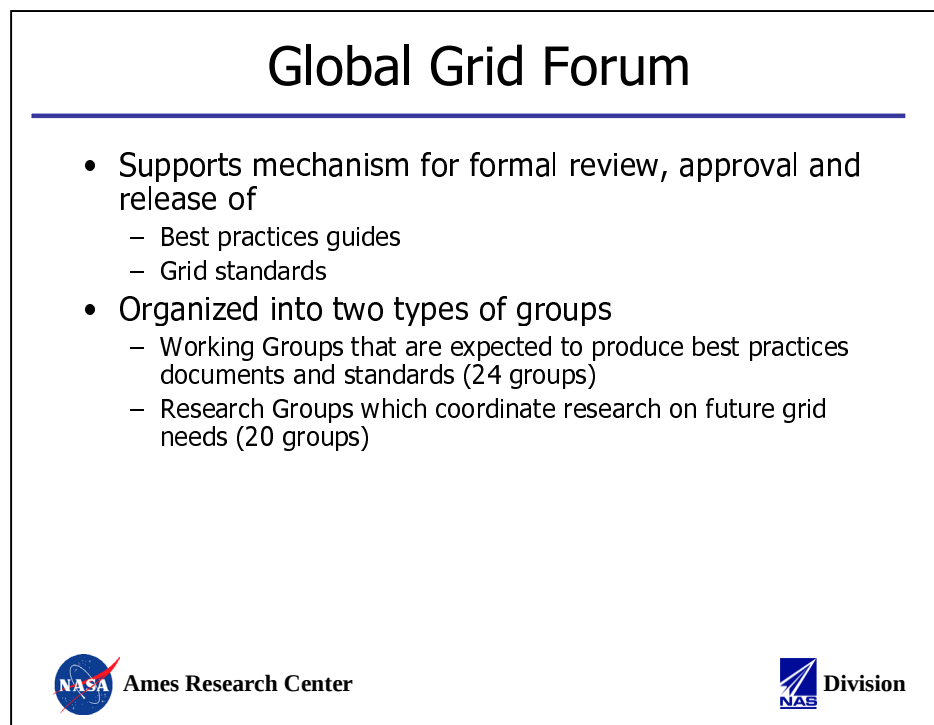


Figure 26

GGF WORKING GROUPS

The slide lists the current GGF Working Groups. Details about each of these groups and the current set of documents and standards on which they are working can be found on the GGF web site at www.ggf.org.

GGF Working Groups

•Grid Checkpoint Recovery	•Discovery and Monitoring Event Description
•New Productivity Initiative	•Network Measurement
•Open Grid Services Architecture	•Grid Information Retrieval
•Open Grid Services Interface	•Previous activities of the Peer to Peer Working Group
•Open Source Software	•Distributed Resource Management Application API
•Data Access & Integration Services	•Grid Economic Services Architecture
•GridFTP	•Grid Resource Allocation Agreement Protocol
•Authorization Frameworks and Mechanisms	•OGSA Resource Usage Service
•Certificate Authority Ops	•Scheduling Attributes
•Grid Certificate Policy	•Scheduling Dictionary
•Grid Security Infrastructure	•Usage Record
•Open Grid Service Architecture Security	
•CIM based Grid Schema	

 **Ames Research Center**

 **Division**

Figure 27

GGF RESEARCH GROUPS

The slide lists the current GGF Research Groups. Details about each of these groups can be found on the GGF web site at www.ggf.org.

GGF Research Groups

• Advanced Collaborative Environments	• Data Replication
• Advanced Programming Models	•Data Transport
• Applications and Test Beds	• Grid Benchmarking
• Grid Computing Environments	• Relational Grid Information Services
• Grid User Services	•Appliance Aggregation (
• Life Sciences Grid	•OGSA-P2P-Security
• Production Grid Management	• Grid High-Performance Networking
• Accounting Models	• Persistent Archives
• Grid Protocol Architecture	• Site Authentication, Authorization, and Accounting Requirements
• Semantic Grid	
• Service Management Frameworks	

 **Ames Research Center**

 **Division**

Figure 28

WHY IS THE GLOBAL GRID FORUM IMPORTANT

The primary reason that the GGF is important is that it will result in grid standards and grid standards will encourage commercial companies to make grid products that satisfy these standards. Standard based products should be more marketable than products that do not satisfy standards.

In addition the GGF provides an arena for various application-specific requirements to be injected into the international grid community. Currently there are a number of application-specific research groups at GGF that may, as the need is found, develop application-specific standards or influence other standards work to address needs unique to a particular application area.

Why is the Global Grid Forum Important

- It will result in grid standards
 - It will encourage commercial products since there will be standards which the products can meet
 - Products that meet accepted standards should be more marketable
- It provides a forum to get application-specific requirements injected into the grid development efforts

 Ames Research Center Division

Figure 29

REUSABLE COMPONENTS FOR GRID COMPUTING PORTALS

Marlon Pierce
Indiana University
Bloomington, IL

Reusable Components for Grid Computing Portals

Marlon Pierce
Community Grids Lab
Indiana University



Figure 1

Grids Today and Tomorrow

- Grid software enables loosely coupled, globally distributed computing
 - “Virtual Organizations”.
- What does that really mean?
 - Specific services such as global authentication, resource allocation management, aggregated information services
 - Centered around a few wire protocols and service implementations
- What’s next? **Open Grid Service Architecture**
 - Use XML (WSDL) to provide a service definition language.
 - Extend WSDL to support metadata about services.



Figure 2

What Is Missing?

- Grids are designed to enable Virtual Organizations.
 - Inter-organizational collaboration
- But we must also support the Real User
 - Provide access to the Grid from any computer (or anywhere).
 - Provide user interfaces to Grid services.
 - Provide customizable front ends that contain the service front ends.
- Grid Computing Environments
 - Browser-based Web portals

Figure 3

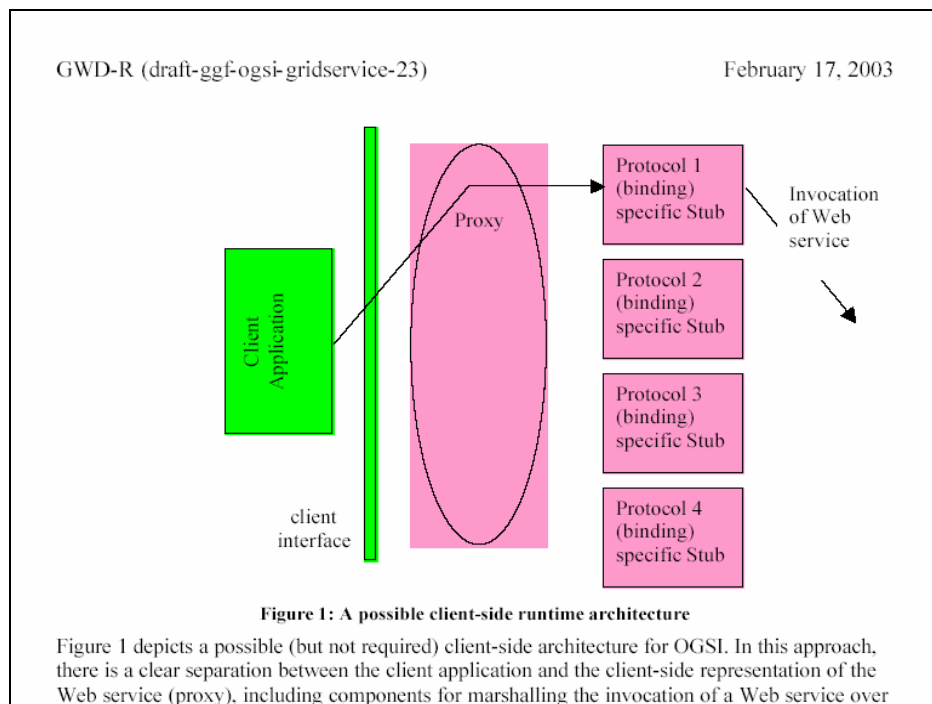


Figure 4

Grid Computing Environments



- Organizations setting up Grids have seen the value of developing user environments, or *Grid Computing Environments*.
 - 28 articles in November-December 2002 issue of *Concurrency and Computation: Practice and Experience*.
 - IPG Launchpad, HotPage, Alliance Portal, and others
- World-wide development community interacts through the GCE research group in the Global Grid Forum.
 - G. Fox (IU), D. Gannon (IU), and M. Thomas (TACC) co-chair.
- Grid portal technology is coming of age
 - Reusability of components
 - Common frameworks

Figure 5

Example GCE: Gateway Portal



- Developed for DOD supercomputing centers (ARL and ASC MSRCs).
- Support source-restricted (commercial or otherwise) applications
 - Ansys, Abaqus, ZNSFlow, Fluent
- Developed to support typical, if simple, high performance computing services
 - Batch script generation, job submission and monitoring, file management and transfer.
 - Do it all securely

Figure 6

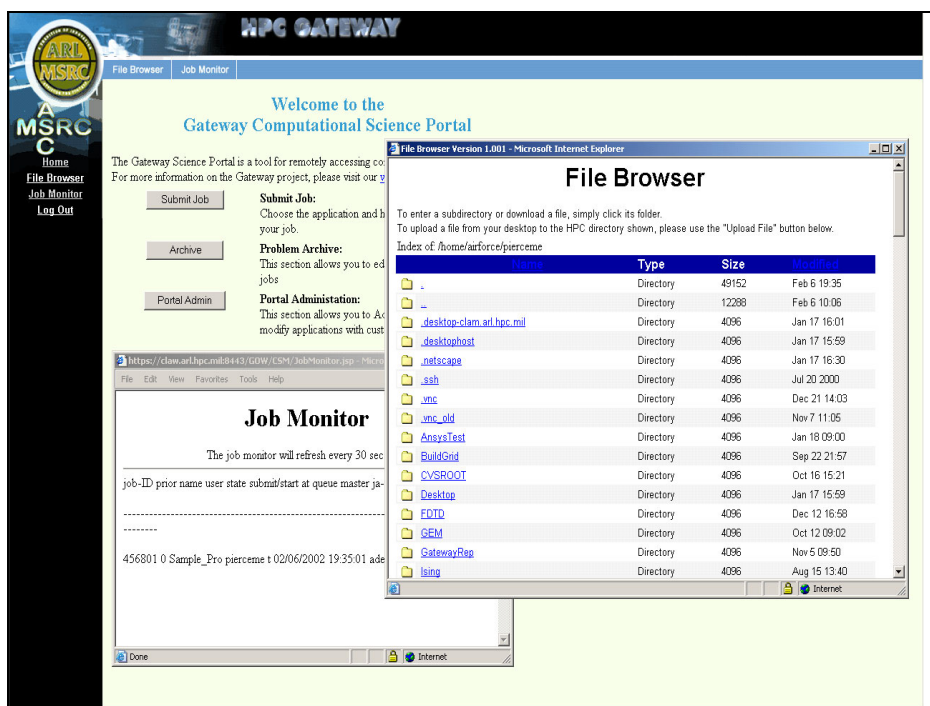
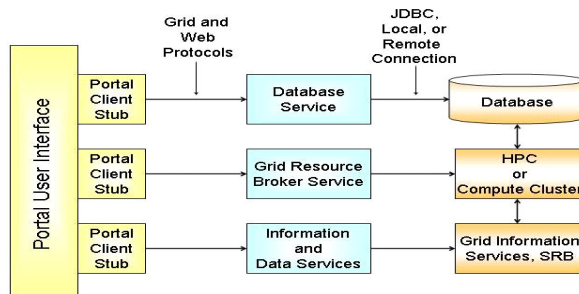


Figure 7

Characteristics of Portals

- Framework contains user interfaces to the services.
- Backend services accessed through service proxies.
- The convergent/emergent architecture is a three tiered model.



The three-tiered architecture is a standard for accessing Grid and other services.

Figure 8

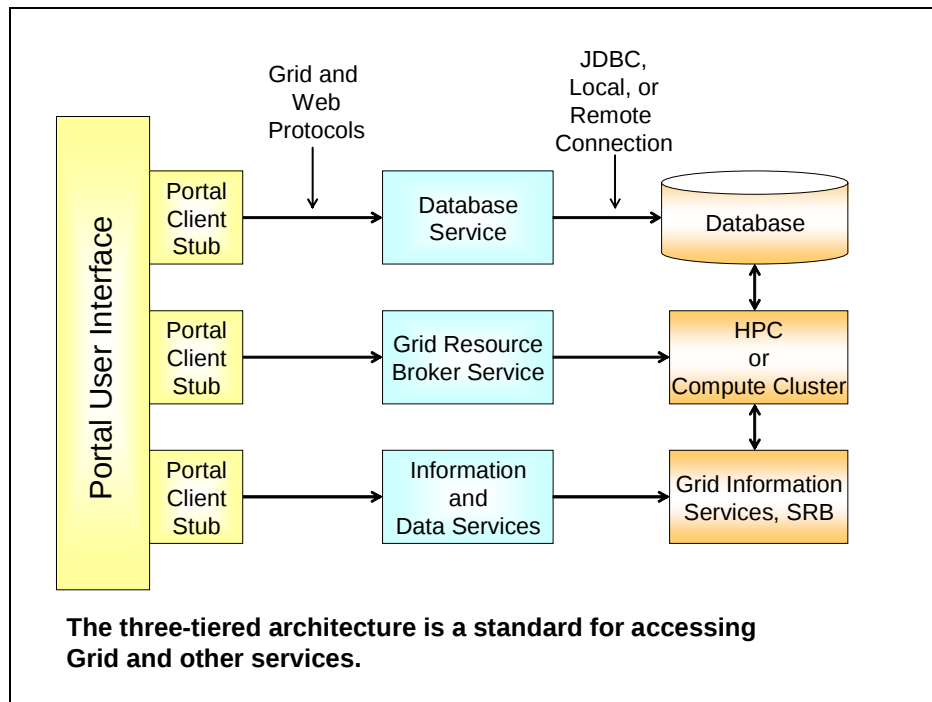


Figure 9

Sharing Portal Services



- Given that everyone builds essentially around the same architecture
 - How do I build a client to interact with someone else's services?
 - How do I build a compatible service implementation?
 - How can I take someone else's end-to-end solution and plug it into my portal.
 - How do I avoid reinventing basic services like login, view customization, access restrictions on interfaces.
- To explore possible solutions, we chose to implement a new portal project, QuakeSim, around the **Web services** and **Portlet** models.

Figure 10

QuakeSim Portal

- A number of simulation methods for studying earthquakes are being developed by GEM consortium including:
 - Simplex, Disloc (**JPL**)
 - Virtual California (**UC-Davis**)
 - PARK codes (**Brown**)
- As codes become more robust and accepted, problems emerge:
 - Need to manage information about distributed data sources: multiple databases, sensors, simulated data.
 - Need to organize, manage information about multiple code installation sites.
 - Need to simplify access to data, use of codes, and use of visualization/analysis tools for broad range of users
 - Need to link together
- NASA funded activity to develop **SERVOGrid** Interoperability framework

Figure 11

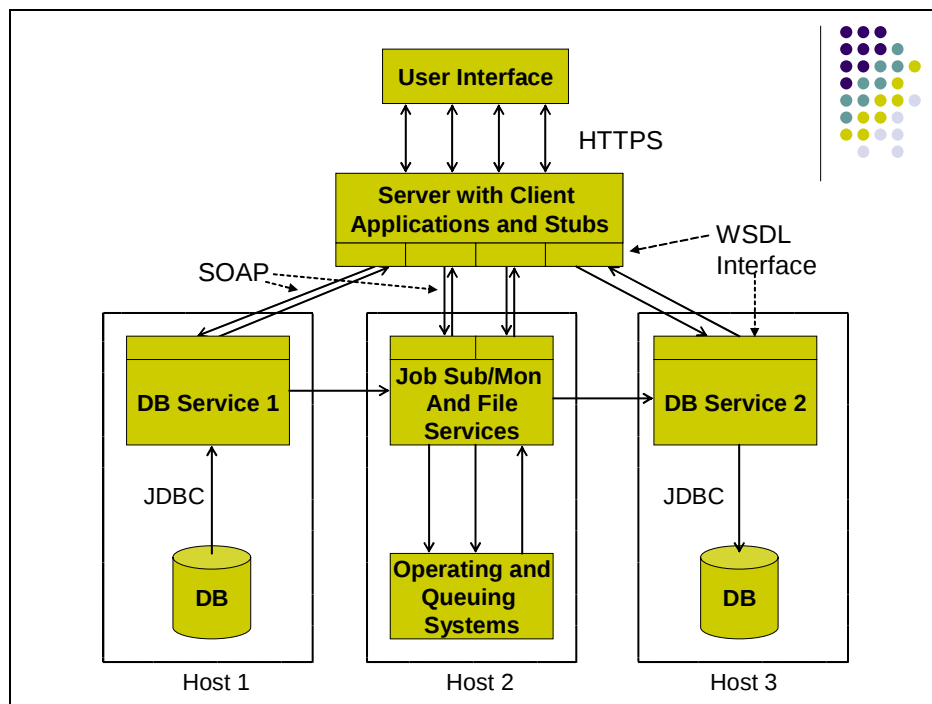


Figure 12

Computing Portal Grid Web Services

- We have built a suite of general purpose Grid Web services for managing distributed applications.
- **Core Computing services** define general purpose functions:
 - Ex: job submission, file transfer, job monitoring, management of jobs and results
 - Described as a **GridShell** as plays same role to Grid that Shell does for UNIX on a single machine
- **Application Grid Web services** include metadata about applications.
 - Built on top of core services.
 - Original application NOT changed
- We have developed a **toolkit** that allows one to convert general software packages into Grid Web Services and manage application collections

Figure 13

Application Grid Web Services



- AGWS are designed to make scientific applications (i.e. earthquake modeling codes) into Grid Resources
- AGWS services are described by two XML Schemas:
 - **Abstract descriptors** describe application options. Used by the application developer to deploy his/her service into the portal.
 - **Instance descriptors** describe particular user choices and archive them for later browsing and resubmission.

Figure 14

User Application Selection and Submission

GEMPortal | [Solar.uts.indiana.edu](#) | [Noahsark.ucs.indiana.edu](#) | [Application Admin](#)

GEM Portal

Code Selection

Please select a code and host machine from the following list of applications. When you have made your choice, click the "Make Selection" button at the bottom of the page.

Disloc

- ☐ solar.uts.indiana.edu
- ☐ grids.ucs.indiana.edu
- ☐ Test2

Simplex

- ☐ solar.uts.indiana.edu
- ☐ grids.ucs.indiana.edu
- ☐ noahsark.ucs.indiana.edu
- ☐ Test2

VC_STRESS_GREEN

VC_SG_COMPRESS

VC_INIT_SER

VC_SER

Welcome to JitSpeed -> [Home](#)
[Customize](#) | [Logout](#) | [Edit Account \(gateway\)](#)
 Server date: Dec 4, 2002 5:11:47 PM

PBS Script Generator

Please provide the following information needed to generate the PBS queue script that will be run on solar.uts.indiana.edu.

Amount of Memory:

Job Name:

WallTime (hh:mm:ss):

Number of CPUs:

Email: ☐ When job begins Email ☐ When job ends Email ☐ When job aborts Email

The code you have selected takes 1 input file. See the code documentation for details.

Input File:

The application generates 1 output file. Please specify the full path name of the directory on the HPC server where you would like this file to be placed.

Output File:

Figure 15

Administer Grid Portal

GEMPortal | [Solar.uts.indiana.edu](#) | [Noahsark.ucs.indiana.edu](#) | [Application Admin](#)

GEM Application Admin Portal

Application Update

Please provide the following information needed to update an application.

Please select a code from the following list of applications:

- ☐ Disloc
- ☐ Simplex
- ☐ VC_STRESS_GREEN
- ☐ VC_SG_COMPRESS
- ☐ VC_INIT_SER
- ☐ VC_SER

Code: Disloc

Input parameters, set 1

Input Field Handle: GEM input

Input Description: The code you have selected takes 1 input file. See the code documentation for details.

Input Mechanism: C-Style Arguments

Output parameters, set 1

Output Field Handle: GEM output

Output Description: The application generates 1 output file. Please specify the full path name of the directory on the HPC server where you would like this file to be placed.

Output Mechanism: C-Style Arguments

Host Computers and Queues:

Host Name	Host IP	Queue Type	Memory Directive	Job Name Directive	Number of Nodes Directive	Walltime Directive	Email Options Directive
solar.uts.indiana.edu	129.79.5.104	PBS	#PBS-l mem=	#PBS-lN	#PBS-lncpus=	#PBS-l walltime=	
grids.ucs.indiana.edu	156.56.103.5	CSH					

Figure 16

Portlets for Reusable Portal Components



- What we found was that groups did not really want to use common interfaces so much as share end-to-end services (user interfaces-client stubs-service implementations).
- Portlets/containers provide a simple way to do this.
- The container implements all portal specific services
 - Manages user customizations, logins, access controls
- Container treats all web content as generic 'portlet' objects.
 - Controls which portlets are displayed and how they are arranged.
- Portlets and containers are implemented in Java
 - Tomcat webapp

Figure 17

The screenshot shows the NCSA Xportlets web application. At the top, there are logos for NCSA, the Globus Project, and Java CoG Kit. A navigation bar contains links to various portlets: Home, Proxy Manager, LDAP Browser, GridFTP Browser, Component Browser, Application Coordinator, XEvents Writer, XEvents Tools, Gram Job Launcher, Newsgroups, Job submission, Job monitoring, TACC, and Collection. Below the navigation bar, there are two main sections. The left section is titled 'xportlets: GridFTP Client' and shows an error message 'Error: proxy null'. The right section is titled 'xportlets: LDAP Browser' and displays the LDAP Server URL 'palomar.extreme.indiana.edu:389' and the current directory 'dc=cs,dc=indiana,dc=edu'. Below these sections, there is a 'Direction:' section with instructions on how to use the application. At the bottom, there is a form to fill out parameters for retrieving a proxy from the MyProxy server, including fields for Hostname, Port, Username, Password, and Lifetime.

Figure 18

Value of the Portlet Approach



- With portlets, we have a common infrastructure for managing content.
 - I don't have to reinvent login, user customization services.
 - But I may choose to add my own service implementation in a well defined way.
- Content (and service user interfaces) are added in a well defined way
 - Edit an xml registry file.

Figure 19

Portlet Implementations



- Several groups (IU, TACC, NCSA, UMich) are using Jetspeed
 - Open source portlet implementation from Jakarta
- We extend it to
 - Add custom services for message boards, chats, etc.
 - Develop specific portlets to Grid services (like GridFTP).
 - Build general purpose portlets to support needs of Grid service interfaces
 - Session state conversations, multipage content, security
 - Bridge to legacy JSP and non-Java Web interfaces

Figure 20

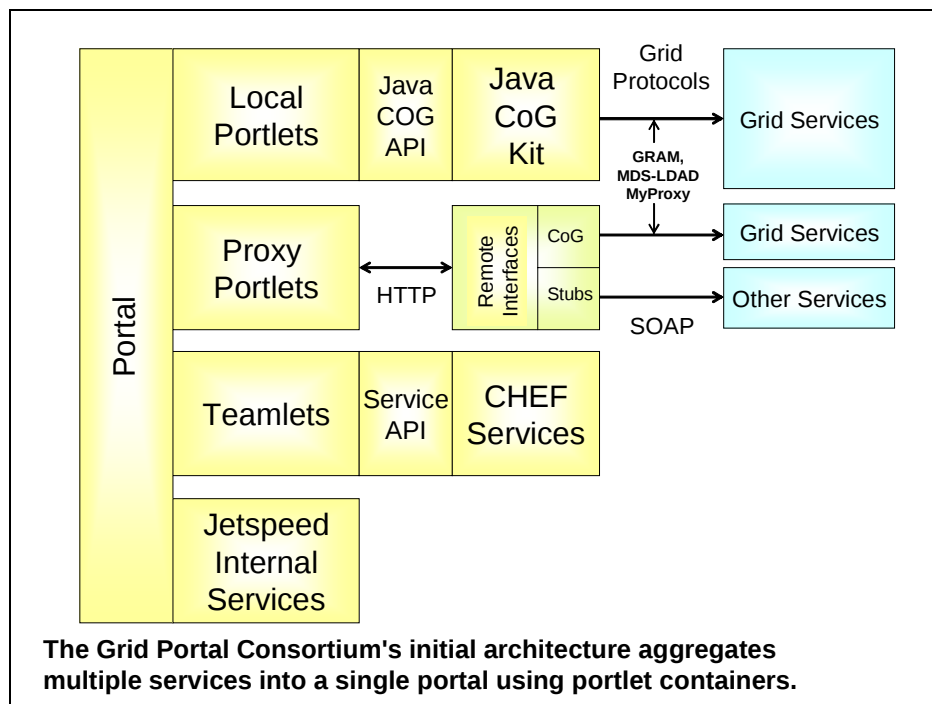


Figure 21

Portlet Longevity

- Portlets have become popular in commercial enterprise servers
- The portlet API is being standardized through the Java Community Process.
 - Participants include IBM, Oracle, BEA, and others.
- We anticipate or will contribute to building the open source reference implementation of the standard.

Figure 22

Portlets and Portal Stacks

- User interfaces to Portal services (Code Submission, Job Monitoring, File Management for Host X) are all **managed as portlets**.
- Users, administrators can customize their portal interfaces to just precisely the services they want.

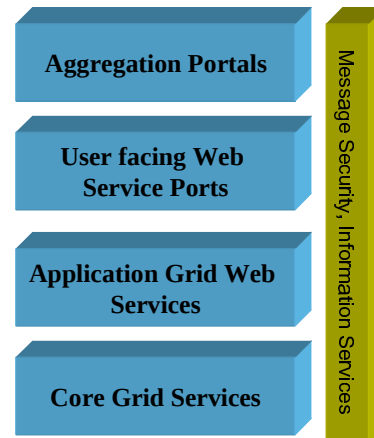


Figure 23

Future Developments

- User interfaces and services need to get much more sophisticated, intelligent.
 - Case-based reasoning interface for Earthquake simulation codes.
 - More standard collaboration services as portlets
 - Whiteboards, chat interfaces
 - Ubiquitous access in a standard fashion
- Portlet repositories to allow community sharing of reusable components.

Figure 24

More Information



- My Email: marpierc@indiana.edu
- Gateway homepage: www.gatewayportal.org
- More publications:
<http://ptlportal.ucs.indiana.edu>.

Figure 25

**Research by Federal Agencies That Will Affect
Future Computing Paradigms for Aerospace**

David Nelson
National Coordination Office for Information Technology Research and Development
Arlington, VA

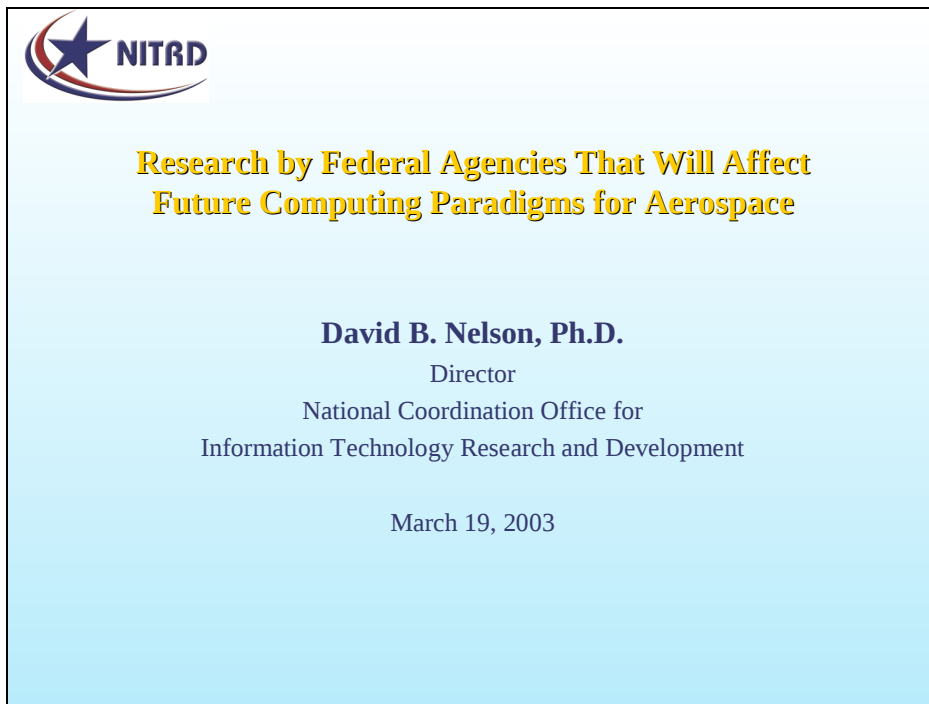


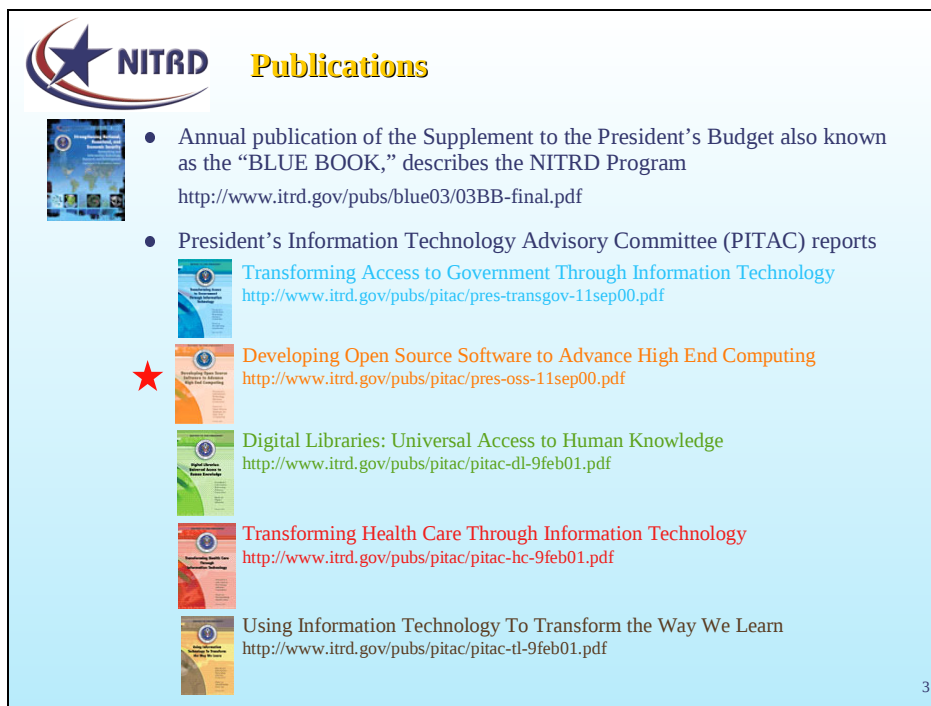
Figure 1



Figure 2

PUBLICATIONS

Publications of the President's Information Technology Advisory Committee include "Developing Open Source Software to Advance High End Computing," that was handed out at the workshop. Open source software is software for which the human-readable source code is made widely available, either as public domain software, or copyrighted with a license that requires source code to be made available. Open source software is an important emerging factor that will affect future aero-space computing. Discussion of open source would be a useful topic for this meeting, but time did not allow its inclusion in this talk.



The slide features the NITRD logo (a blue star with a red swoosh) and the title "Publications" in yellow. It lists several publications with corresponding thumbnail images. The third item, "Developing Open Source Software to Advance High End Computing", is highlighted with a red star. Each item includes a title and a URL.

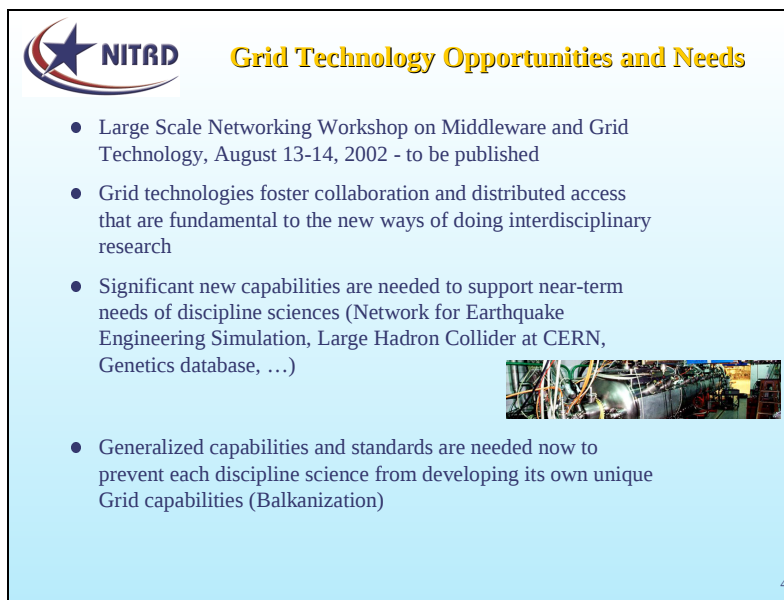
-  • Annual publication of the Supplement to the President's Budget also known as the "BLUE BOOK," describes the NITRD Program
<http://www.itrd.gov/pubs/blue03/03BB-final.pdf>
- President's Information Technology Advisory Committee (PITAC) reports
 -  [Transforming Access to Government Through Information Technology](http://www.itrd.gov/pubs/pitac/pres-transgov-11sep00.pdf)
<http://www.itrd.gov/pubs/pitac/pres-transgov-11sep00.pdf>
 -  ★ [Developing Open Source Software to Advance High End Computing](http://www.itrd.gov/pubs/pitac/pres-oss-11sep00.pdf)
<http://www.itrd.gov/pubs/pitac/pres-oss-11sep00.pdf>
 -  [Digital Libraries: Universal Access to Human Knowledge](http://www.itrd.gov/pubs/pitac/pitac-dl-9feb01.pdf)
<http://www.itrd.gov/pubs/pitac/pitac-dl-9feb01.pdf>
 -  [Transforming Health Care Through Information Technology](http://www.itrd.gov/pubs/pitac/pitac-hc-9feb01.pdf)
<http://www.itrd.gov/pubs/pitac/pitac-hc-9feb01.pdf>
 -  [Using Information Technology To Transform the Way We Learn](http://www.itrd.gov/pubs/pitac/pitac-tl-9feb01.pdf)
<http://www.itrd.gov/pubs/pitac/pitac-tl-9feb01.pdf>

3

Figure 3

GRID TECHNOLOGY OPPORTUNITIES AND NEEDS

The workshop on Middleware and Grid Technology, organized by the Large Scale Networking Coordinating Group, produced a report that will be published shortly. Some conclusions of the workshop are presented in these viewgraphs.



NITRD **Grid Technology Opportunities and Needs**

- Large Scale Networking Workshop on Middleware and Grid Technology, August 13-14, 2002 - to be published
- Grid technologies foster collaboration and distributed access that are fundamental to the new ways of doing interdisciplinary research
- Significant new capabilities are needed to support near-term needs of discipline sciences (Network for Earthquake Engineering Simulation, Large Hadron Collider at CERN, Genetics database, ...)
- Generalized capabilities and standards are needed now to prevent each discipline science from developing its own unique Grid capabilities (Balkanization)



4

Figure 4



NITRD **Grid Technology Needs, Concluded**

- Industry is not focused on the longer term research needed to further develop the Grid. Federal research is needed.
- New technical capabilities are needed
 - Testbeds and prototypes for simulations and laboratories
 - Persistent, reliable, high-performance infrastructure
 - Grid economics and accounting
 - Security implementation
 - Standards applying across disciplines and international boundaries
 - Policies for interacting, sharing, and accounting
 - Multidisciplinary, robust, easy-to-use Grid technology and tools

5

Figure 5

GRID COMMUNITES AND APPLICATIONS: HIGH ENERGY PHYSICS PROBLEMS SCALE

Physics data from the Compact Muon Solenoid, a detector on the Large Hadron Collider at CERN in Geneva, Switzerland, will be managed using grid technology. The grid is hierarchical, in that data flow primarily from top to bottom and at each stage the flow rate decreases. The intent of the grid is to provide “seamless interaction” by physicists with each other and with the data.

Similar grid structure, including Open Grid Systems Architecture and the globus toolkit, could be applied to large-scale NASA missions such as the Earth Observing System.

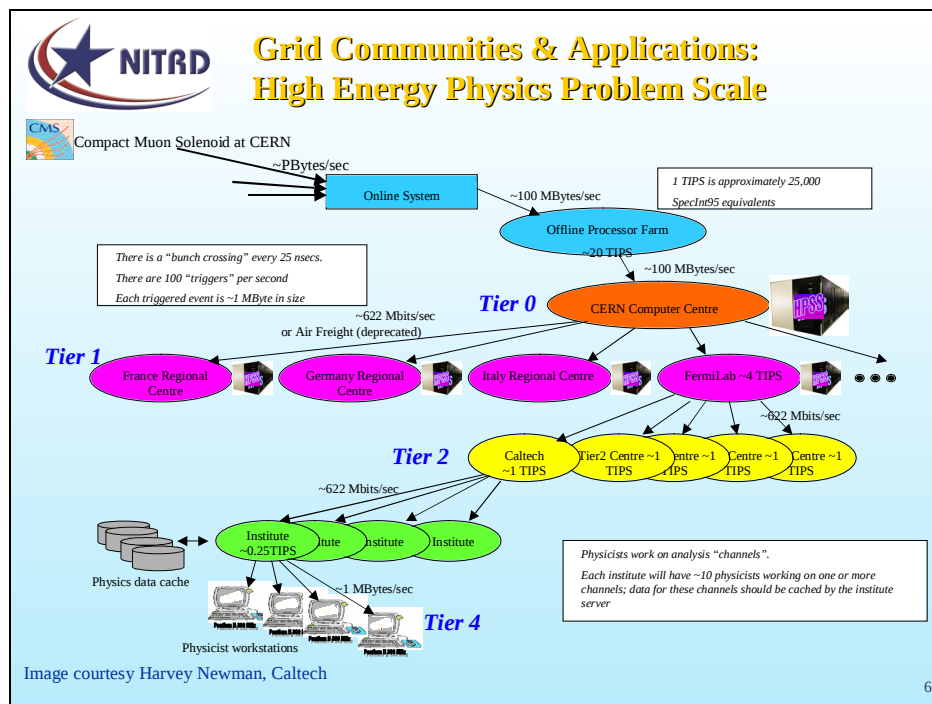


Figure 6

GRID TECHNOLOGY SCENARIO FROM WORKSHOP

The Virtual National Airspace Simulation Environment is a NASA-based scenario from the workshop. The scenario includes dealing with an in-flight emergency that cripples the airplane and requires special pilot responses. The viewgraph lists the grid technology requirements needed for this scenario.



Grid Technology Scenario from Workshop

- Virtual National Airspace Simulation Environment
- Grid Technology Requirements
 - Access to distributed computational resources to support real-time simulations
 - Access to distributed simulation models
 - Access to distributed information resources
 - Real-time access to on-line sensor data, e.g. weather sensors, on-board aircraft sensors
 - Priority for commanding use of resources
 - Security,
 - Reliability, robustness for critical functions
 - Collaboration technology and user interfaces
 - Real-time monitoring and management of Grid tools and resources

7

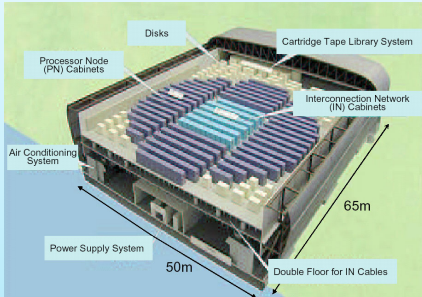
Figure 7



Earth Simulator Has Inspired a New Look at U.S. High End Computing

- Based on the NEC SX architecture, 640 nodes, each node with 8 vector processors (8 GFlop/s peak per processor), 2 ns cycle time, 16GB shared memory
 - Total of 5104 total processors, 40 TFlop/s peak, and 10 TB memory
- Has a single stage crossbar switch(1800 miles of cable) 83,000 copper cables, 16GB/s cross section bandwidth
- 700 TB disk space
- 1.6 PB mass store
- Area of computer = 4 tennis courts, 3 floors

8



Source: <http://www.es.jamstec.go.jp/esc/eng/outline/outline02.html>

Figure 8

PERFORMANCE MEASURES OF SELECTED TOP COMPUTERS

This viewgraph presents three performance measures for the computers at the top of the Top 500 supercomputers (www.top500.org.) The performance measures include R-peak, the peak theoretical performance of the computer measured in Giga-Flops/second, R-max, the best performance on the Linpack program, also measured in Giga-Flops/second, and the Stream Triad benchmark, that gives the memory access rate for the calculation $C(I) = A(I) + Q*B(I)$ for very large vectors, measured in Giga-Bytes/second.

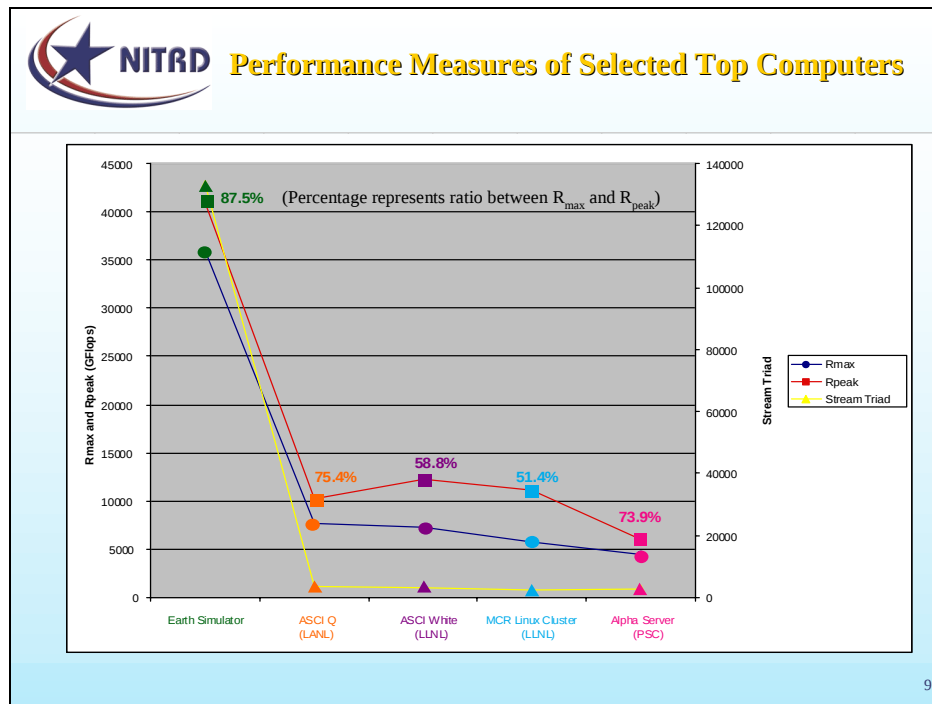


Figure 9



Several Federal Agencies Have Recently Examined High End Computing Needs

- They are mostly using COTS-based HEC
- Most expect COTS to be acceptable in near term, however:
 - Time-to-solution becoming too long
 - Too hard to program; too hard to optimize
 - Coordinated improvements are needed in hardware, software, and application algorithms
 - Rapidly escalating demand on HEC facilities
- Some important applications/algorithms are not amenable to COTS-based HEC
 - Primarily due to non-local memory reference e.g., long vectors requiring gather-scatter operations

10

Figure 10



Examples of Applications for Which COTS May be Unsuitable

- Hypersonic air-breathing propulsion
 - Needs high memory-to-CPU bandwidth for multi-disciplinary analysis
- Reusable Launch Vehicle Design
 - Needs high memory-to-CPU bandwidth
- Protein Folding
 - Poorly parallelizable
- Cryptoanalysis
 - Needs fast flat-memory model
- Climate data assimilation
 - Part of problem not easily parallelizable, needs high memory-to-CPU bandwidth

11

Figure 11



Agency Conclusions

- Further progress in HEC will require balanced, coordinated effort in
 - Research, development, and engineering of new HEC architectures and systems
 - Procurement of new COTS and custom systems
 - Better software (systems, middleware, and applications)
 - Better domain science (mathematics and algorithms)
- HEC is a decreasing part of the technical computing marketplace.
- COTS-based HEC is largely based on technologies developed for low- and mid-range markets (SMP nodes, low bandwidth interconnects).
- Market pressure may result in future COTS systems being less responsive to HEC needs.
- Federal funding of highest-performing HEC, including development of new systems, may be required.

12

Figure 12



High End Computing Revitalization Task Force (HECRTF) Charge

- Rationale: High End Computing (HEC) increasingly critical
- HECRTF coordinated through National Science and Technology Council (NSTC)
- To develop a plan that can guide future Federal HEC investments
- Plan will lay out an overall strategy for these investments
- Seek wide participation by Federal agencies developing or using HEC
- Final report to be completed by August 2003, in time to be an input to FY 2005 budget

13

Figure 13

High Productivity Computing Systems

Robert Graybill
Defense Advanced Research Projects Agency
Arlington, VA

INTRODUCTION

High performance computing is at a critical juncture. Over the past three decades, this important technology area has provided crucial superior computational capability for many important national security applications. Unfortunately, current trends in commercial high performance computing, future complementary metal oxide semiconductor (CMOS) technology challenges and emerging threats are creating technology gaps that threaten continued U.S. superiority in important national security applications.

As reported in recent DoD studies, there is a national security requirement for peta-scale high productivity computing systems. Without government R&D and participation, high-end computing will be available only through commodity manufacturers primarily focused on mass-market consumer and business needs. This solution would be ineffective for important national security applications.

Improving system performance is no longer sufficient to increase system productivity. DARPA'S High Productivity Computing Systems (HPCS) program must also improve system programmability, portability, and robustness. HPCS is pursuing the research and development of balanced, economically viable high productivity computing system solutions for the national security and industrial user communities.

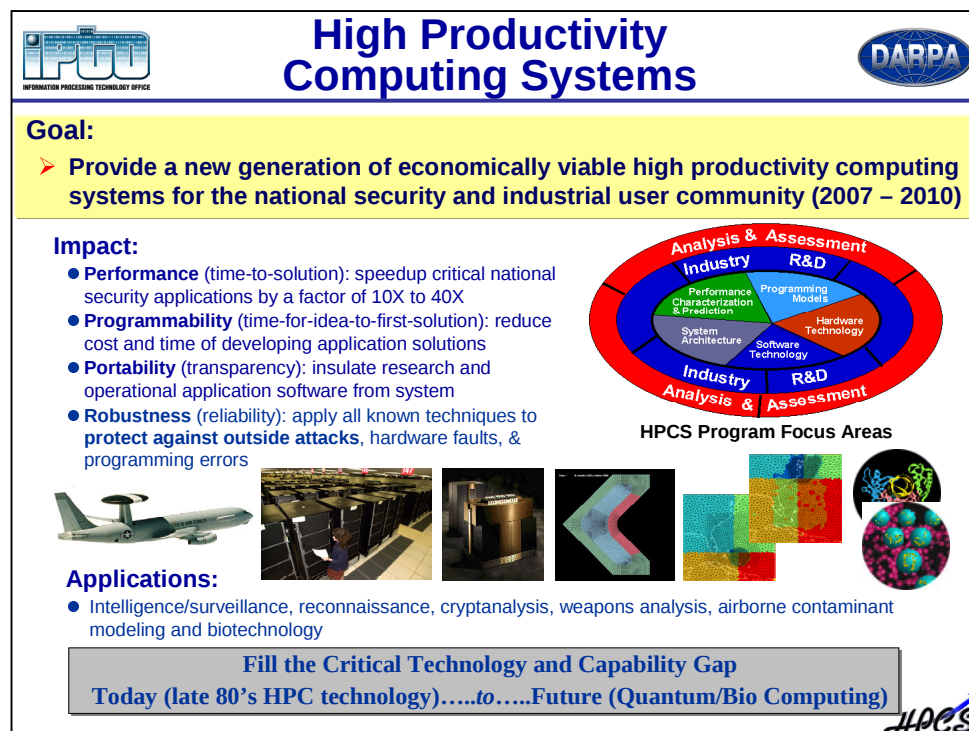


Figure 1

VISION

Today's high-end systems tend to fall into one of two domains: the vector supercomputer domain or the commodity high performance computer domain. Foreign computer vendors dominate the vector domain with Cray as the sole domestic supplier. A majority of the tera-scale computing installations in the United States consist of commodity HPCs.

The High Productivity Computing Systems program will bridge the gap between the late-80's based technology of today's High Performance Computers and the promise of quantum computing for the Department of Defense. DARPA's challenge is to develop a broad spectrum of innovative technologies and architectures integrated into a **balanced** total system solution by the end of this decade.

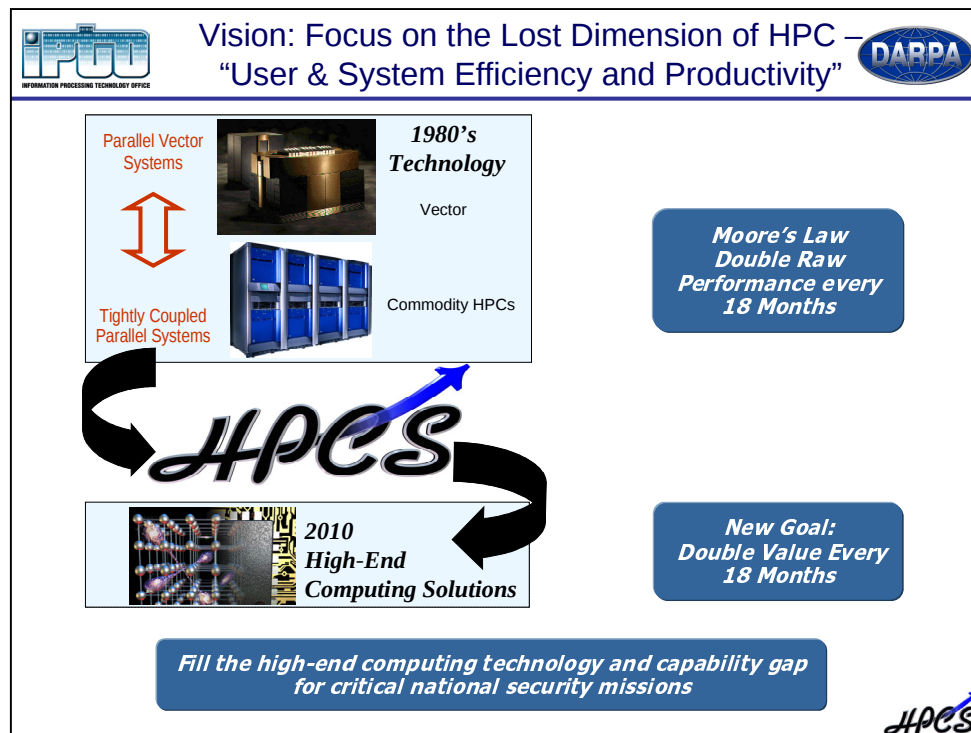


Figure 2

SCHEDULE

To achieve the aggressive goal of revolutionary HPCS solutions by the end of this decade, three top-level program phases have been identified to address the challenges of scalable vector and commodity HPC solutions for today and tomorrow. The three phases are concept study, research and development, and full-scale development. The one-year Phase I industry concept study, completing in June 2003, will provide critical technology assessments, develop revolutionary HPCS concept solutions, and supply new productivity metrics necessary to develop a new class of high-end computers by the end of this decade.

The second phase of the HPCS program is a three-year research and development effort that will perform focused R&D and risk reduction engineering activities. These pursuits will result in a series of system design reviews, preliminary design reviews and risk reduction prototypes and demonstrations. The technical challenges and promising solutions identified during the concept study will be explored and prototyped by a full complement of commercial industry, university, and research laboratory researchers.

Phase III, full-scale development, will be led by commercial industry. This phase will last four years and complete the detailed design, fabrication, integration and demonstration of the full-scale HPCS pilots.

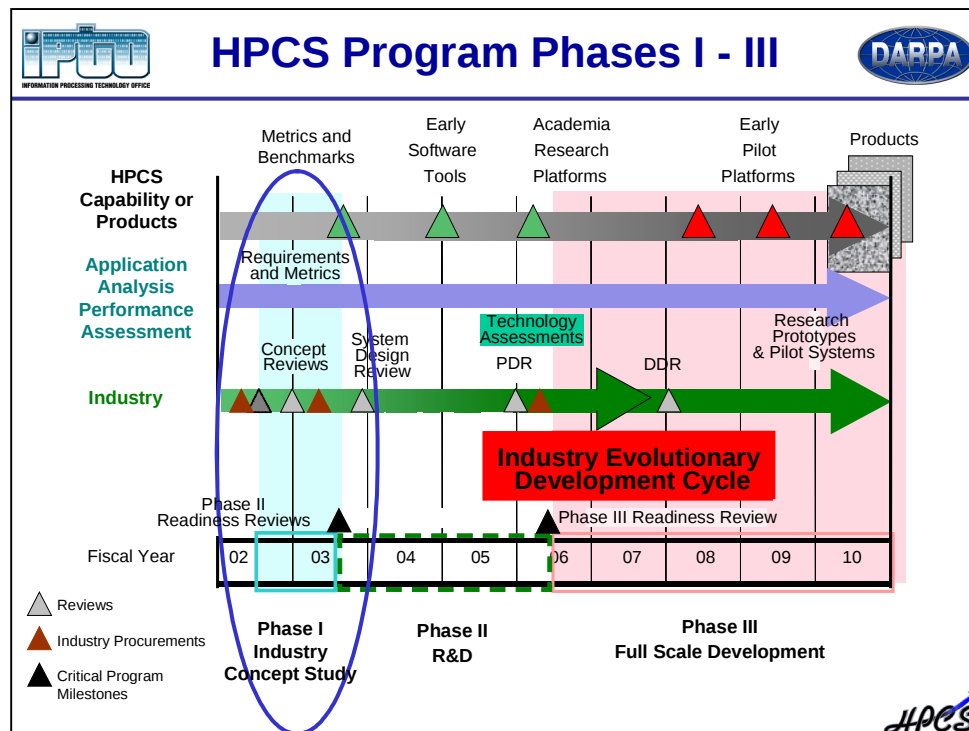


Figure 3

HPCS TEAMS

Phase 1 HPCS vendors are listed below. HPCS concept study awards were made to industry teams led by IBM, Cray, SGI, Sun, and HP. Across these industry teams at least 20 different universities are represented.

Throughout all three phases of the program, application analysis and performance assessment activities will be carried out. Some work is done by the HPCS vendors, for their own benefit. Much of the work is being done by national labs, universities and other organizations for the benefit of the entire HPCS program. These organizations make up the Applications Analysis and Performance Assessment Team. The team is being led by MITRE and MIT/Lincoln Laboratory.

**HPCS Phase I Industry Teams**

Industry:

	Cray, Inc. (Burton Smith)
	Hewlett-Packard Company (Kathy Wheeler)
	International Business Machines Corporation (Mootaz Elnozahy)
	Silicon Graphics, Inc. (Steve Miller)
	Sun Microsystems, Inc. (Jeff Rulifson)

Application Analysis/Performance Assessment Team:





Figure 4

APPLICATION ANALYSIS AND PERFORMANCE ASSESSMENT

The Application Analysis and Performance Assessment Team is studying those mission areas identified as having inadequate available computational resources and is coordinating with the HPCS Mission Partners to identify challenge applications that will serve as the requirements drivers for HPCS. The challenge application selection process started with inputs from the DDR&E and Integrated High-End Computing (IHEC) Mission Analysis studies, which identified areas where deficiencies in the present computing capabilities exist that affect mission performance. Consultations with HPCS Mission Partners generated lists of actual operational and research codes and an understanding of the partners' software development processes and system utilization patterns. The team has identified full-scale applications, compact applications and kernels that represent the mission partners' needs and supplied them to the HPCS Phase I vendors. The team is working to profile these applications and to characterize the underlying requirements in parallel with the HPCS vendors.

The Application Analysis and Performance Assessment Team has also worked with the HPCS Phase I vendors on development of HPCS productivity metrics, and a framework that puts them into a concise context for HPCS.

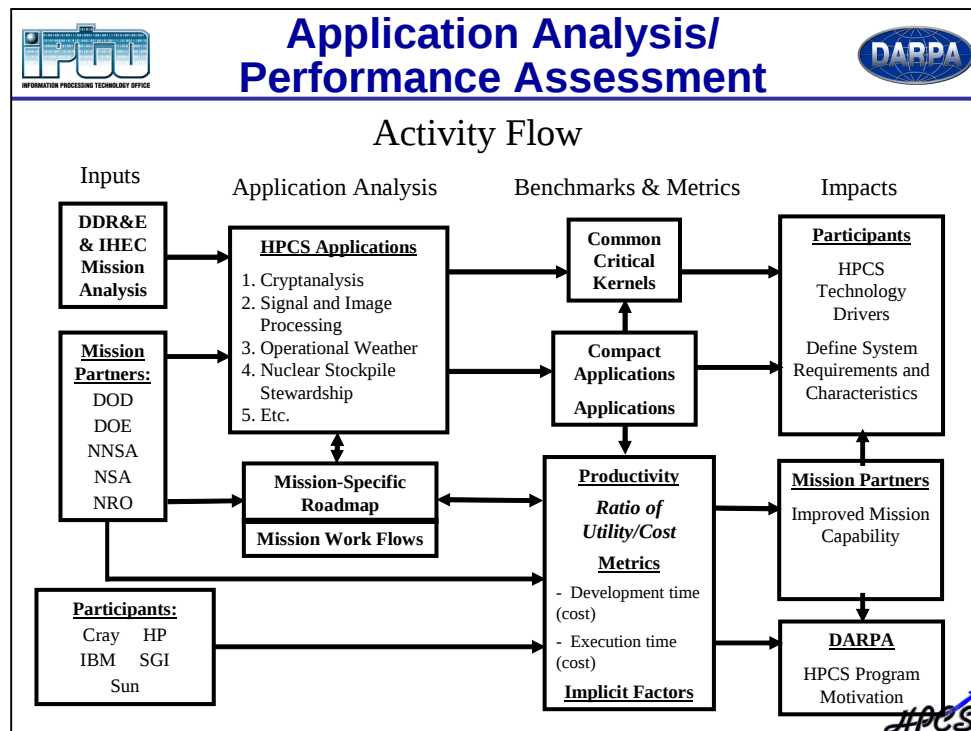


Figure 5

APPLICATIONS

As reported in recent DoD studies, there is a national security *requirement* for high productivity computing systems. Without government R&D and participation, high-end computing will be available only through commodity manufacturers primarily focused on mass-market consumer and business needs. The HPCS program will significantly contribute to DoD and industry information superiority in at least the applications areas colored in red and blue on the chart. The HPCS program will create and supply new systems and software tools that will lead to increased productivity of the applications used to solve these critical problems.

The HPCS mission areas highlighted in red and blue were chosen from two studies of national security computing needs. The DDR&E study performed by the Office of the Secretary of Defense focused on the national security requirements for high-end computers. The Integrated High-End Computing (IHEC) Mission Analysis performed at the request of Congress was a much broader study exploring the requirements, key technologies, proposed long implementation/organization strategy, and funding projections.

	Application Focus Selection	
DDR&E Study • Operational weather and ocean forecasting • Planning activities for dispersion of airborne/waterborne contaminants • Cryptanalysis • Intelligence, surveillance, reconnaissance • Improved armor design • Engineering design of large aircraft, ship and structures • National missile defense • Test and evaluation • Weapon (warheads and penetrators) • Survivability/stealth design	IHEC Study • Comprehensive Aerospace Vehicle Design • Signals Intelligence (Crypt) • Signals Intelligence (Graph) • Operational Weather/Ocean Forecasting • Stealthy Ship Design • Nuclear Weapons Stockpile Stewardship • Signal and Image Processing • Army Future Combat Systems • Electromagnetic Weapons Development • Geospatial Intelligence • Threat Weapon Systems Characterization	

Figure 6

WORKFLOWS

In conducting interviews with HPCS Mission Partners during Phase I, the Application Analysis and Performance Assessment Team found that three general workflows are representative of the Partners' operations and needs. **Workflows** identify how Mission Partners use HPCs—they describe the iterative processes of software development and system utilization and define mission partners' priorities.

The workflows that characterize HPCS missions are lone researcher, enterprise development and production/operations. For each class of user, the total time to solution is strongly dependent upon the coupling that exists between execution time and development time. The diagrams on the left represent a high-level view of the operational workflow, while the diagrams on the right represent the software development workflows. For example, the first row of the chart depicts the workflow of the “Lone Researcher.” His or her goal is to rapidly understand and solve a domain-specific problem. The overall execution cycle is characterized by rapid iterations between the development of new hypotheses or theories and testing those theories computationally. The development model is characterized by rapid prototyping. This is very different from the production/operations workflow in row three. Here the goal is to create a fielded system that will rapidly process external inputs to provide actionable data to decision makers. The overall execution cycle is driven by real-time considerations. The development cycle consists of both an initial development of the system and a maintenance cycle once it is fielded.

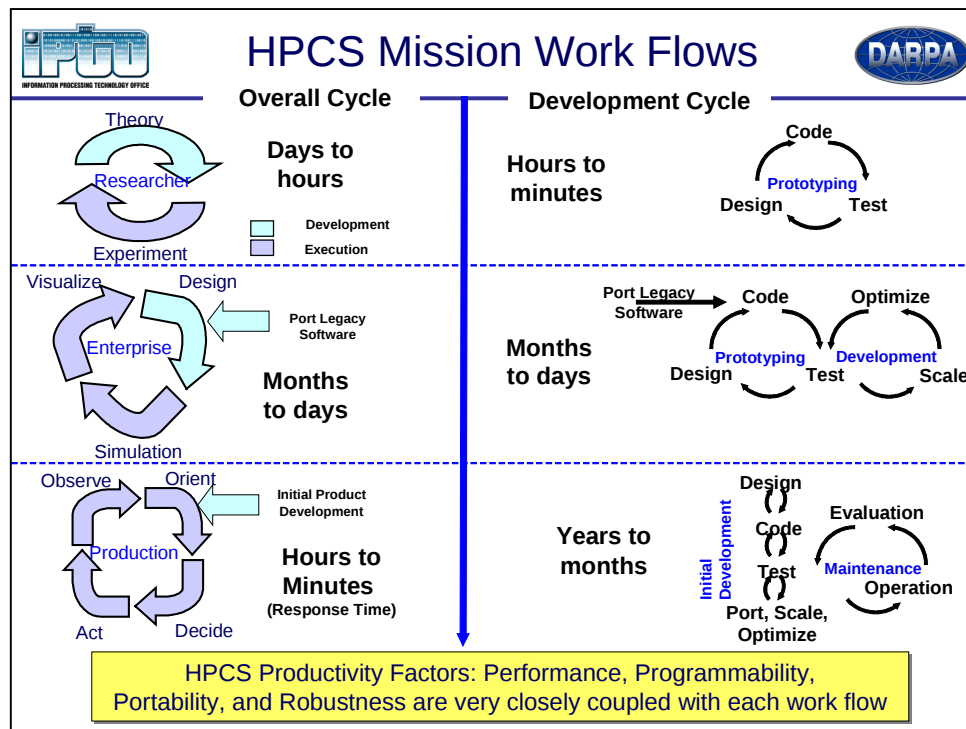


Figure 7

ASSESSMENT FRAMEWORK

Common metrics (such as peak floating-point operations per second) are insufficient for understanding and assessing system capabilities. The Application Analysis and Performance Assessment Team has worked with mission partners and HPCS vendors to develop an appropriate assessment framework. The initial framework shown consists of: Productivity Metrics (e.g. development time and execution time); System Parameters (e.g. bandwidth, flops/cycle, size, power, lines-of-code); Workflows, Benchmarks and Systems models. The system parameters and benchmarks are depicted as inputs into an actual or modeled HPCS system and generate productivity metrics. The productivity metrics are depicted as inputs into mission workflow models, which can be used to determine the productivity (or value) of a particular system for a particular mission. Workflows provide insights on how the various mission partners will evaluate HPCS systems.

Implicit Productivity Factors such as performance, programmability, portability and robustness are attributes of the both the system and the workflow and reflect the system capabilities and the needs of the mission.

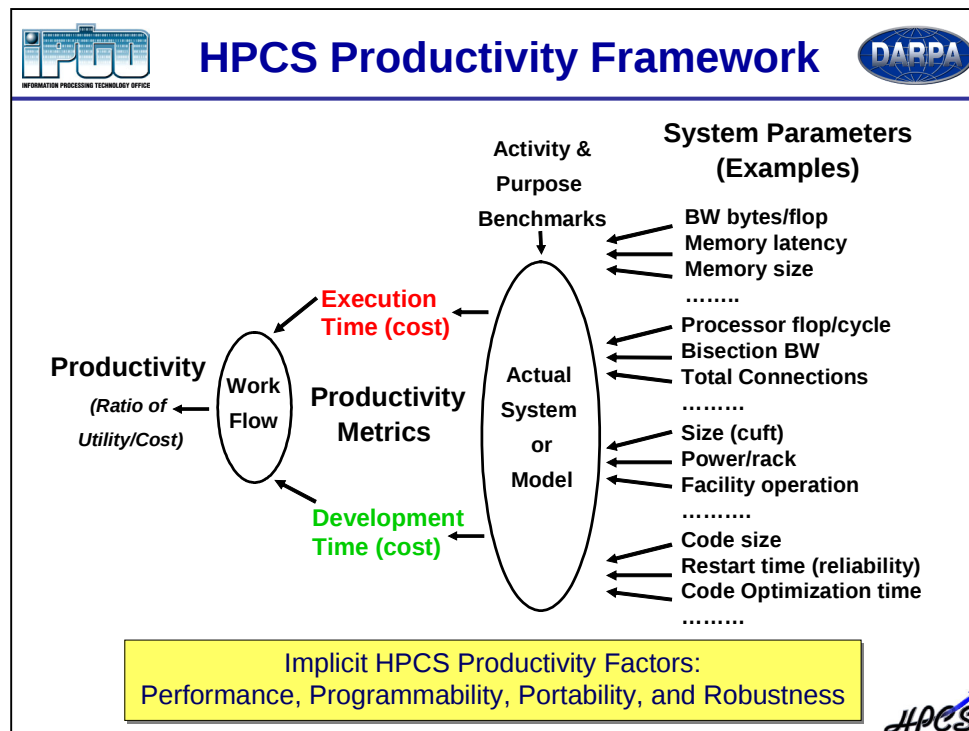
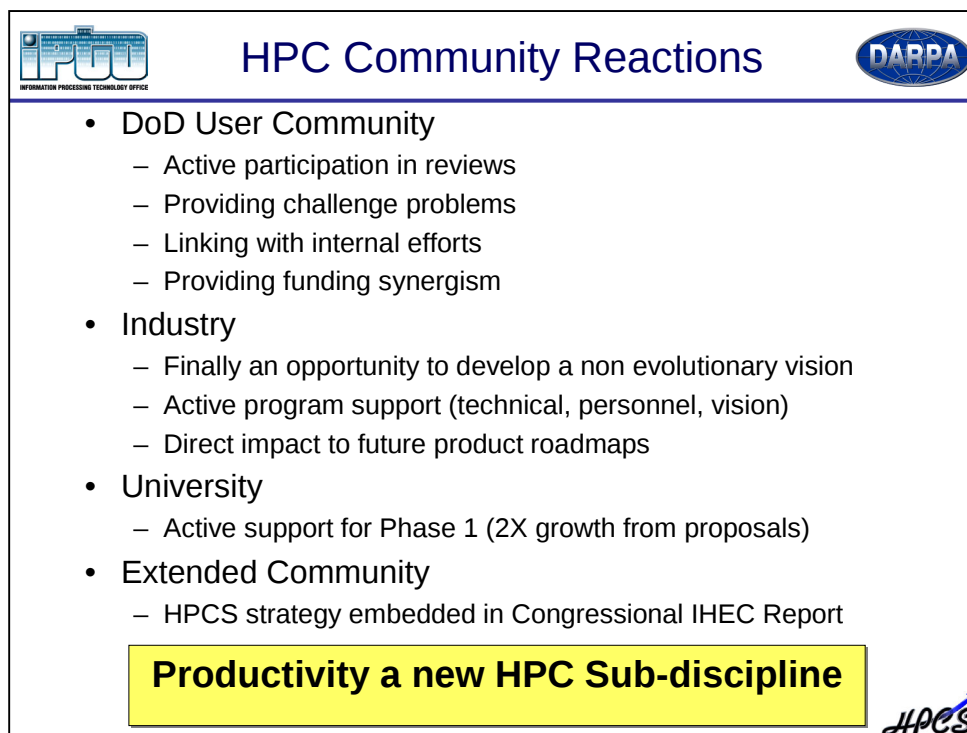


Figure 8

SUMMARY

The HPCS program has received very positive response from the vendor, government, and university communities. HPCS represents the first comprehensive high-end computing program since the early 90's. The focus on productivity or the ability to easily program highly parallel systems with high sustained performance across a spectrum of computing applications represents not only a significant challenge but an opportunity to fill a major gap in realizable parallel computing. HPCS provides the vendors with an incentive to break out of the current evolutionary computing development paradigm by exploring new innovative technologies, architectures, and programming techniques. The very active and synergistic participation of the DoD users, universities, and vendors in all phases of this program is beginning to pay off. HPCS is laying the foundation for future larger scale programs such as the one proposed in the IHEC Congressional Study Report.



HPC Community Reactions

- DoD User Community
 - Active participation in reviews
 - Providing challenge problems
 - Linking with internal efforts
 - Providing funding synergism
- Industry
 - Finally an opportunity to develop a non evolutionary vision
 - Active program support (technical, personnel, vision)
 - Direct impact to future product roadmaps
- University
 - Active support for Phase 1 (2X growth from proposals)
- Extended Community
 - HPCS strategy embedded in Congressional IHEC Report

Productivity a new HPC Sub-discipline

Figure 9

**BUILDING A COLLABORATIVE BRIDGE – TECHNOLOGY
RESEARCH, EDUCATION, and COMMERCIALIZATION CENTER**

Jonas Talandis
National Center for Supercomputing Applications
University of Illinois Urbana/Champaign
Chicago, IL

BUILDING A COLLABORATIVE BRIDGE – TECHNOLOGY RESEARCH, EDUCATION, and COMMERCIALIZATION CENTER

The Technology Research, Education, Commercialization Center (TRECC) is a technology center located in Dupage County, IL, west of Chicago. TRECC is sponsored by the United States Office of Naval Research.



Figure 1

TRECC MISSION

The mission of TRECC is to accelerate the development of innovative ideas, to develop new education applications and learning systems, to demonstrate on-the-horizon information technologies, and to incubate start-up technology businesses.

TRECC Mission

- **Showcase Technology Research**
- **Develop Educational Applications and Learning Systems**
- **Technology Transfer and Commercialization of Emerging Technologies**





Figure 2

WHO ARE WE?

The Technology Research, Education and Commercialization Center, or TRECC, is a University of Illinois at Urbana-Champaign (UIUC) program, managed by the National Center for Supercomputing Applications (NCSA), sponsored by the Office of Naval Research (ONR).

Who Are We?

- **The U of I and NCSA**
 - Program Management, Technology Research, and Continuing Education
- **Office of Naval Research**
 - Funding and Direction
- **Battelle Memorial Institute**
 - Commercialization and Small Business Support





Figure 3

WHAT DO WE DO?

TRECC sponsors development of and showcases technologies developed by our partners in the National Computational Science Alliance (Alliance) of which NCSA is the leading-edge site. Battelle Memorial Institute (Battelle) of Columbus, OH is the sub-contractor responsible for Small Business Assistance and Client Services. They match private-sector industry with appropriate government technologies, resources or interests. UIUC provides Education, Training and Learning Research in the form of Business Education, Continuing Ed, directed Academic Outreach and e-Learning programs.

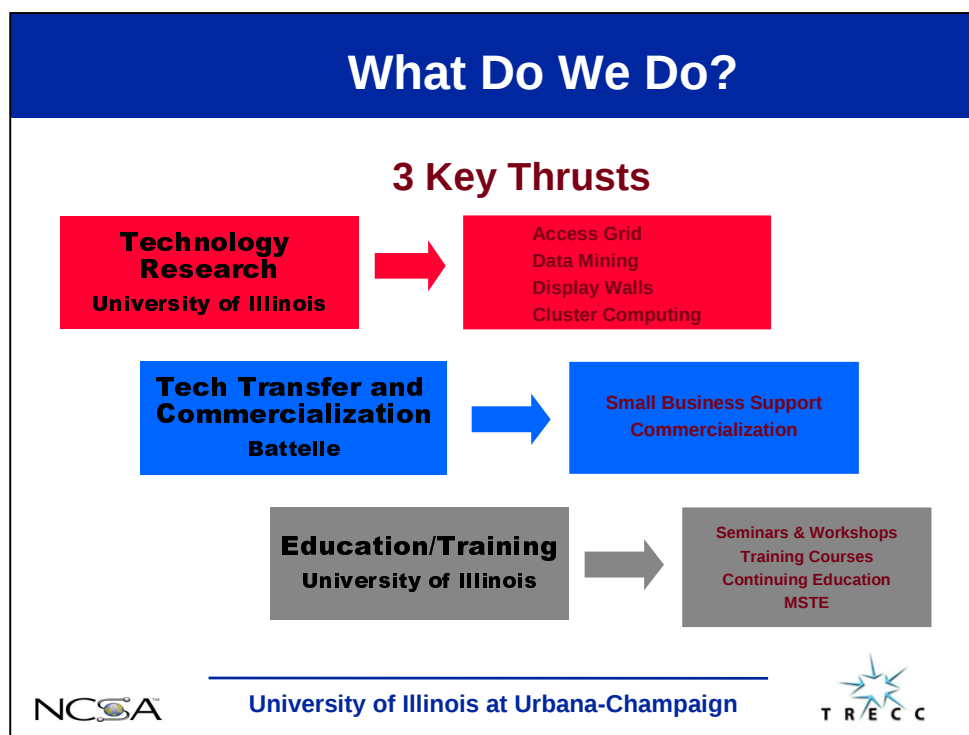


Figure 4

TECHNOLOGY RESEARCH/DEMONSTRATION

TRECC exposes government and businesses to leading edge infrastructure technologies. Our Grid presence provides the community accessibility to resources and serves as a application test-bed for grid-related applications. TRECC's collaborative back-bone is the Access Grid, among other video tele-conferencing (vtc) and immersive technologies, in which we assist our clients in their deployment of.

Technology Research/Demonstration

- **Leading Edge Infrastructure Technologies**
 - **Grid Facilities**
 - Community Portal to National Technology Grids
 - Scientific, Engineering and Educational Applications to Validate the Grid's Commercial Relevance.
 - **Collaborative Technologies**
 - Access Grid Node and other Tele-Immersive
 - Partner Deployment Assistance
 - **Deployable Demonstrations**
 - Distributed Cluster Computing Architectures
 - Advanced Display Technologies
 - Data Mining and Information Visualization Frameworks
 - Collaboration Frameworks for Information Sharing



Figure 5

TRECC EDUCATION

TRECC is an education and training hub to prepare the workforce of today and tomorrow, along with the science and technology teachers of tomorrow. The MSTE program works with the Technology Center of Dupage in curriculum development for teachers that interest students in math and science studies, tools they will need to succeed. TRECC offers workshops and training for businesses and technical professionals. The knowledge center has four key components. A database, collaborative conferencing tools, a knowledge exchange utility and e-Learning environment.

TRECC Education

- **Education and Training Hub**
 - **Math and Science Teacher Education (MSTE)**
 - Vocational Outreach
 - **Continuing Education**
 - Business and Entrepreneurial Training
 - IT and Technical
 - **Knowledge Center**
 - Knowledge Base
 - Collaboration Space
 - Knowledge Exchange
 - e Learning Environment



Figure 6

TECH TRANSFER/COMMERCIALIZATION

Battelle is responsible for TRECC's tech transfer and client service activities. They work to match available government technologies with suitable businesses or can work the other way, bringing private-sector technology to government interests. Small business assistance is provided. A client's technology is assessed for innovative value and marketability. Business plans are outlined and/or reviewed. Funding and opportunity databases are searched for available matches. Clients are assisted throughout the proposal process, in preparing paperwork and preparing conduits for progress.

Tech Transfer/Commercialization

- **Identify Suitable DoD Technologies**
 - Liaison with Tech, Processes and Systems
 - Match with Partners
- **Client Services**
 - Technology Assessment
 - Market Opportunities
 - Partnership Support
 - Information Services
 - SBIR/STTR/BAA Database Searches
 - DOD Funding Opportunities
 - Proposal Support



Figure 7

DUPAGE AIRPORT AUTHORITY

TRECC is located on the third floor of the Flight Center, Dupage Airport, West Chicago, IL. Aviators visiting TRECC can scrape their wingtips on the building when visiting. DPA is the 3rd busiest airport in Illinois, and 11th in the Great Lakes region.



Figure 8

ECONOMIC DEVELOPMENT

TRECC is the result of earmarked federal funds in combination with local and state authorities. TRECC is the first increment in the development of the Dupage Research Park, a 1000+ acre development on airport property. Directors of the research park board represent the Dupage Airport Authority, Dupage County, and the University of Illinois. The research park is expected to include an entrepreneurial technology incubator and satellite university campus(es). Neighboring towns have teamed with Fermi National Laboratory, which borders the site, to expand and offer broadband services to their communities. The development of Global Distributed Technology Centers promises regional sites with facilities enabled for crisis management and security activities.

Economic Development

- **Dupage Research Park**
 - State appropriated \$34M
 - \$5M Released
- **Local, State and Federal Initiatives**
 - Tech Incubator
 - Broadband Services
- **Global Centers**
 - Crisis Management
 - Homeland Security





Figure 9

WHY A CENTER?

Global Centers help get academics and the DoD together and make it easy for investors, partners and all involved to understand a technology's value. Kicking the tires of a nascent technology is important as is effective n-way communications of data, teamwork and relationships.

Why a Center?

- **Timely and Effective Technology Transfer**
 - Get the DoD and Academic Labs Together
 - Understand the “Value Proposition”
- **Get People Together**
 - Meet, Use and Understand the Technology
 - Collaborate – Physically and Virtually
- **Physical Presence**
 - Creates an Identity
 - Transportation and Services are Essential



Figure 10

THE NATIONAL TECHNOLOGY GRID

The Alliance, among other research initiatives are continuously creating new information infrastructures. This map of the National Technology Grid is an example of a continuously changing dynamic.

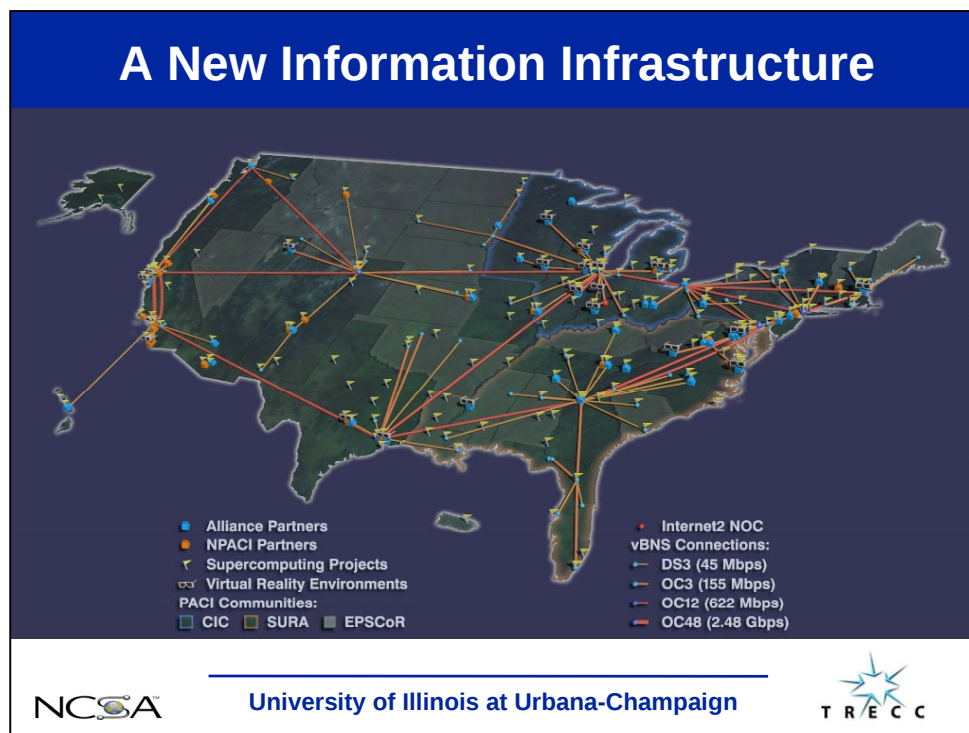


Figure 11

GRID LINKS PEOPLE

The grid provides users with specific or widely distributed resources on a national scale. One vision of penetration is when consumer ‘grid appliances’ such as your phone or an automobile’s information system can access grid resources routinely.

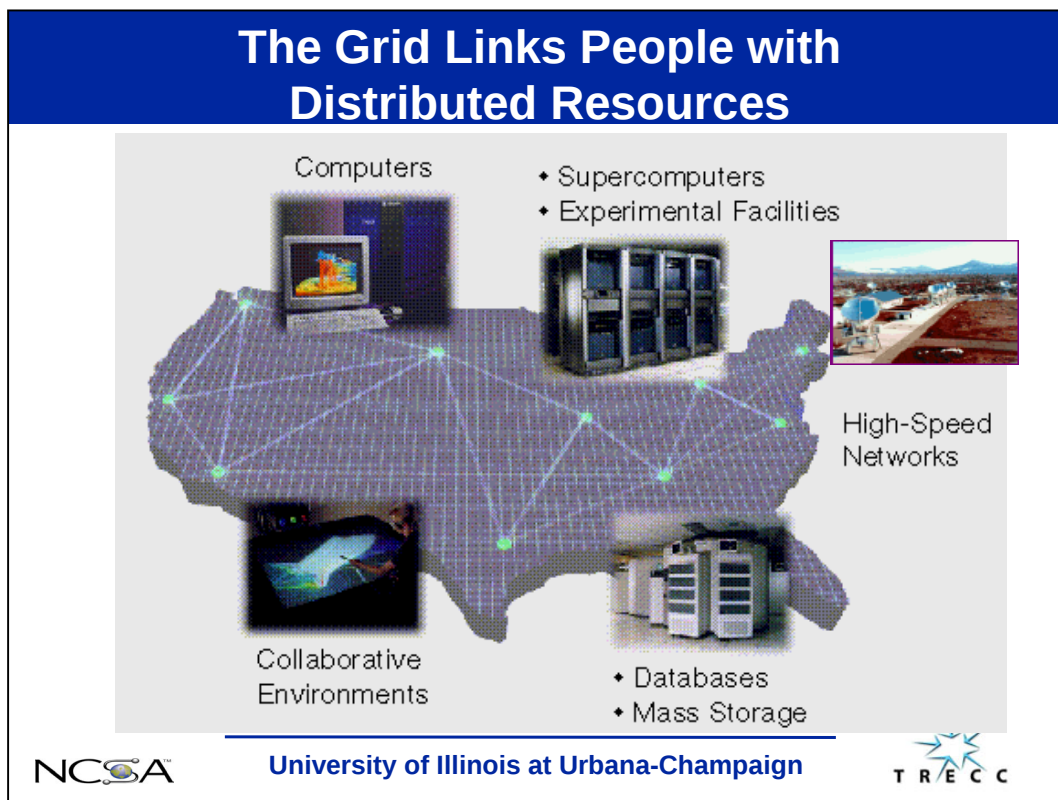


Figure 12

CREATING COLLABORATIVE WORK SPACES

The grid enables people, data, and technology to come together as applications
i.e.. effective productivity tools.

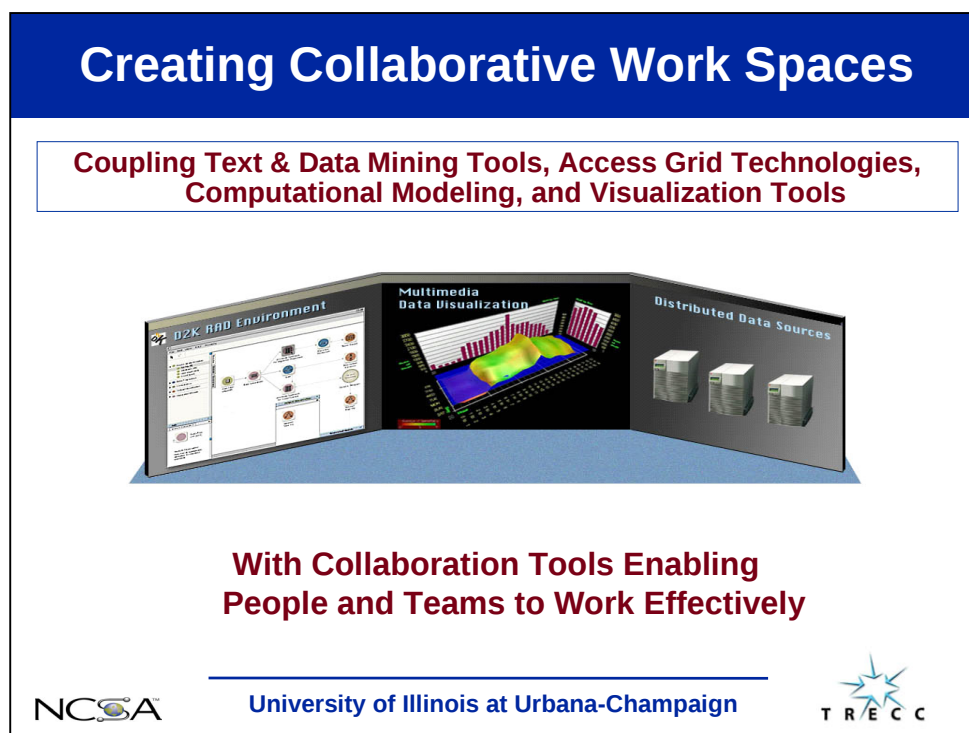


Figure 13

GLOBAL DISTRIBUTED CENTERS MODEL

A Global Center is typified by multiple connectivity options, enabling people to reliably and diversely communicate, thereby collaborate.

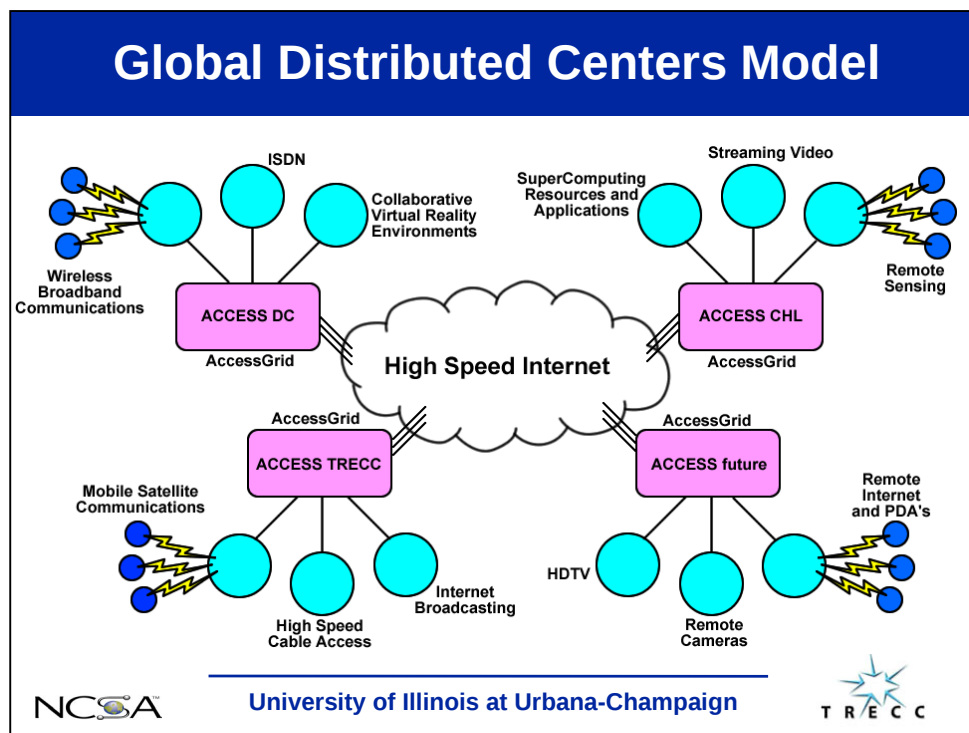


Figure 14

ACCESS - DC

The Alliance Center for Collaboration, Education, Science and Software (ACCESS) located near Washington DC (Arlington, VA) is the prototype center. Additional centers in addition to TRECC exist or are planned in several locations around the globe.



Figure 15

ACCESS and the MSCMC

ACCESS is the host site for the Multi-Sector Crisis Management Center (MSCMC).

ACCESS DC and the MSCMC

ACCESS



Mission: to advance scientific research in computational science through the creation of Global Distributed Centers:

Explore the development and use of advanced technologies

Foster national and international partnerships between academic, government, public, and private sectors

Accelerate technology transfer

MSCMC



Mission: Promote global *Development and Deployment* of advanced IT Strategies and tools for Crisis Management and Emergency Response Communities for all phases (planning, response, mitigation and recovery) including virtual reality environments over all modes of communications.



Figure 16

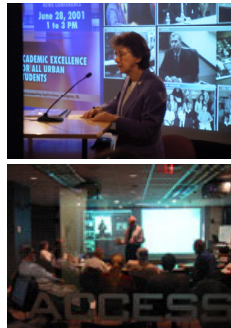
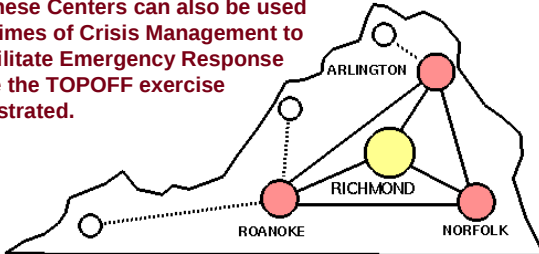
ACCESS and the ACCESS GRID

ACCESS has several initiatives that are regional and national/international in scope. Here is an state-wide example for Virginia. It stresses multi-purpose use of its resources, such as K-12 education assistance in addition to stand-by readiness in cases of crisis management.

ACCESS, MSCMC and the Access Grid Model Example

An ACCESS Centers model offers a Virginia-wide system of Centers connected by a high speed network to deliver collaborative government sessions, meetings and events.

- When not scheduled by the Government for meetings and events these Centers can be used for educational (K-12) activities to bridge the Digital Divide.
- These Centers can also be used in times of Crisis Management to facilitate Emergency Response like the TOPOFF exercise illustrated.



National Center for Supercomputing Applications (NCSA) <http://www.ncsa.uiuc.edu>
Alliance Center for Collaboration Education Science and Software (ACCESS) <http://calder.ncsa.uiuc.edu/ACCESS>
Multi-Sector Crisis Management Consortium (MSCMC) <http://www.msccmc.org> Contact: jtt@ncsa.uiuc.edu

NCSA University of Illinois at Urbana-Champaign **TRECC**

Figure 17

MSCMC ARC NETWORK

The Access Response Network (ARC) outlines global distributed centers in all 50 states, including up to four mobile systems in each state.

MSCMC ARC Network Multi-Mode Access Response Centers

- Centers in all 50 States
- Multiple Mobile Systems
- Multi-Sector Collaboration
- Multiple Use Peace and Prosperity Centers
- Emergency Response
- Economic Development
- Life Long Learning



NCSA University of Illinois at Urbana-Champaign TREC

Figure 18

TOPOFF EXERCISE

The TOPOFF exercise was conducted last year. Agencies leased the ACCESS facility for a period of days as a test and exercise of their emergency response mechanisms.



Figure 19

TRECC BUILD-OUT

TRECC started as raw space in the Dupage Flight Center. The ACCESS consulting architect along with the TRECC program team was responsible for the preliminary design. A local design/build firm provided project architecture and construction services. NCSA staff provided equipment plans and deployment, some of which had to find alternate means of entry into the building.



Figure 20

TRECC FACILITIES

TRECC has five collaboration areas. A Demonstration area accommodates up to 40 people, the Training area, approximately 24. Used together in a common gathering these areas can accommodate approximately 100. Two Conference rooms seat approximately 16 and 8. Studio 6/7 is a demonstration studio for up to 8 people. Staff offices may house 10 full-time employees while guest studios are available for short or longer-term assignment to visitors, as required.



Figure 21

COLLABORATION AREAS

The Demonstration and Training areas feature large format, rear-projected screens, approximately 18' wide x 15' high. Projectors were chosen for their resolution and brightness, including flexibility to accommodate emerging and experimental display technologies as they become available.



Figure 22

CONFERENCE ROOMS

All collaboration areas are well-equipped. Each has full Audio/Visual support, including, microphones, cameras, and appropriate displays. Power outlets and network connections are always close at-hand, including fiber throughout. A Wi-Fi network covers the entire facility, providing wireless connectivity for visitors and staff.



Figure 23

STUDIO 6/7

Studio 6/7 is a demonstration studio, currently showcasing Continuum.



Figure 24

CAPABILITIES

The space, systems and furniture at TRECC are designed for flexible configuration to accommodate various groups/events. A broad range of connectivity options provide access to computational and collaborative resources. A state-of-the-art A/V switching system allows routing of all audio and video signals via touch panel control displays or a via a web interface.

Capabilities

- **Flexible Accommodation**
- **Multiple Broadband Connections**
- **GigE, 100BaseT, WiFi Networking**
- **Facility A/V Support**
- **Access Grid and VTC Enabled**





Figure 25

ACCESS GRID

The Access Grid (AG) is the collaborative backbone of TRECC and is the people part of the grid. AG optimizes group-to-group collaboration, and scales from auditorium environments, to the desktop, down to handhelds.

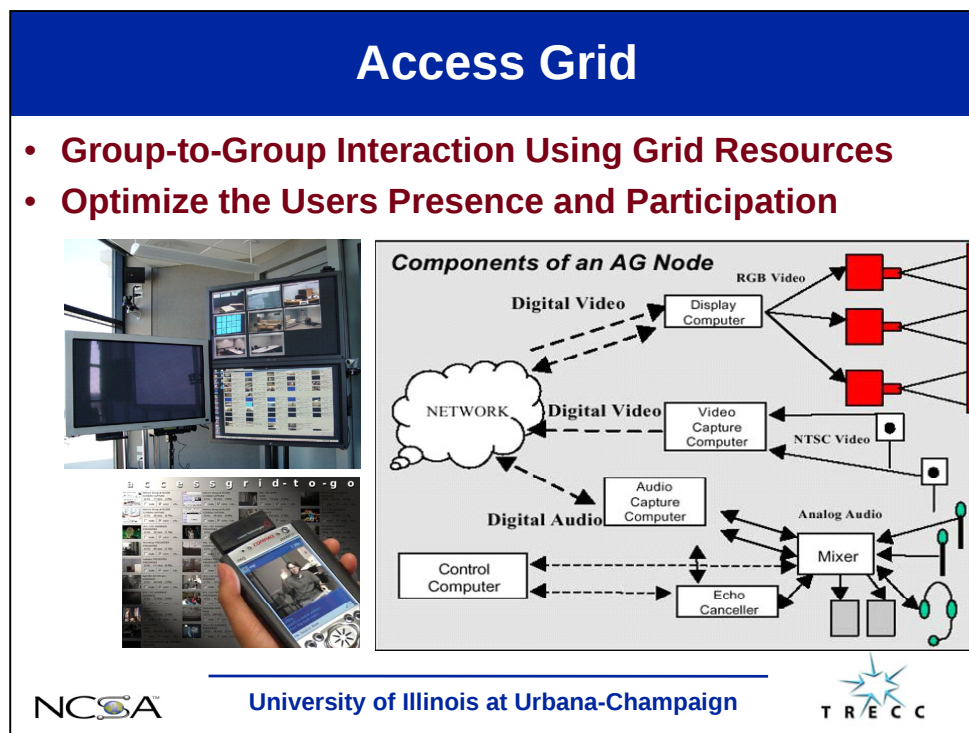


Figure 26

AG- 150+ NODES and COUNTING

The AG display depicted here represents a typical collaborative session. The operator chooses from a pallet of video thumbnails available in a particular virtual venue. They size and position the windows by dragging them across their multiple screen desktop using their mouse. Distributed data and audio streams are included in a full duplex environment.

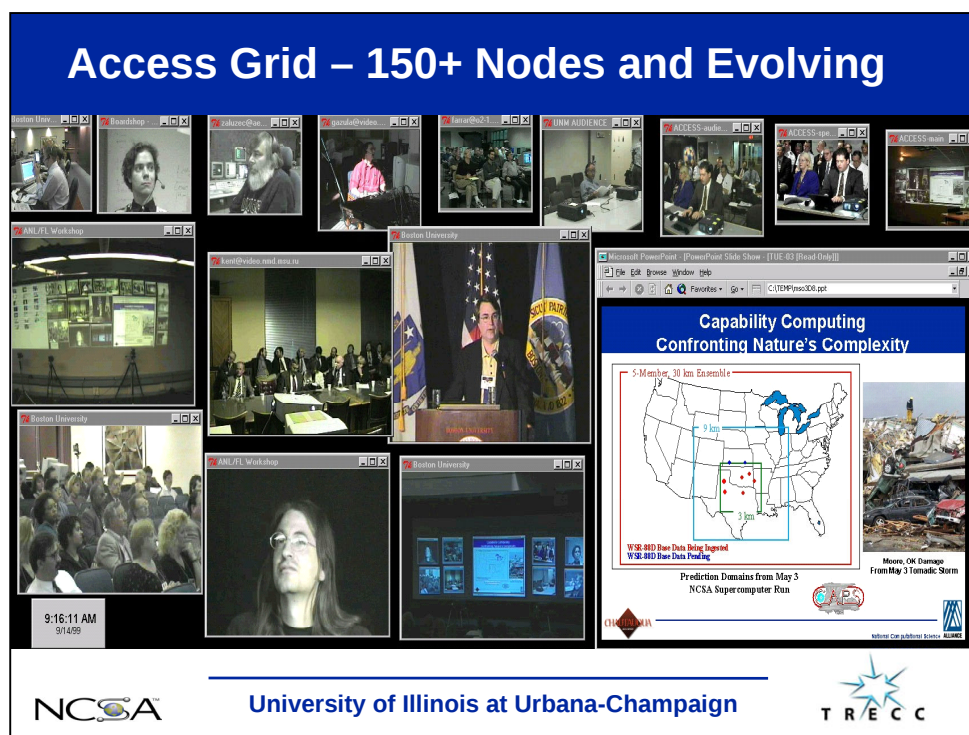


Figure 27

EDUCATION and TRAINING

The training facility accommodates remote and local curricula in a fully supported collaborative environment. A dedicated server supports e-learning application environments.

Education and Training

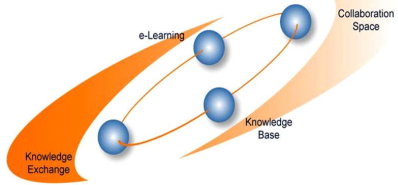
- 22 Seats
- WiFi Laptops
- Training Domain
- AG and VTC Enabled
- Remote or Local Curricula



trecc

Entrepreneurs' Knowledge Center

Home | Knowledge Base | Collaboration Space | Knowledge Exchange | e-Learning



Welcome to the
Entrepreneurs' Knowledge Center

NCSA

University of Illinois at Urbana-Champaign

TRECC

Figure 28

HIGH PERFORMANCE COMPUTING SUPPORT

TRECC's Equipment Room accommodates networking, computational cluster, collaborative and domain servers, with room for expansion.



Figure 29

ADVANCED NETWORKING

With High Performance Computing, you need Advanced Networking. TRECC currently has a dedicated OC-3 (155Mbps) connection to Chicago NAP where we peer with MREN and other research networks. We are currently contracting with SBC for GigaMAN connection to StarLight, on the downtown campus of Northwestern University. Plans are being drawn to expand our connectivity to 10-20 Gigabit in the future.

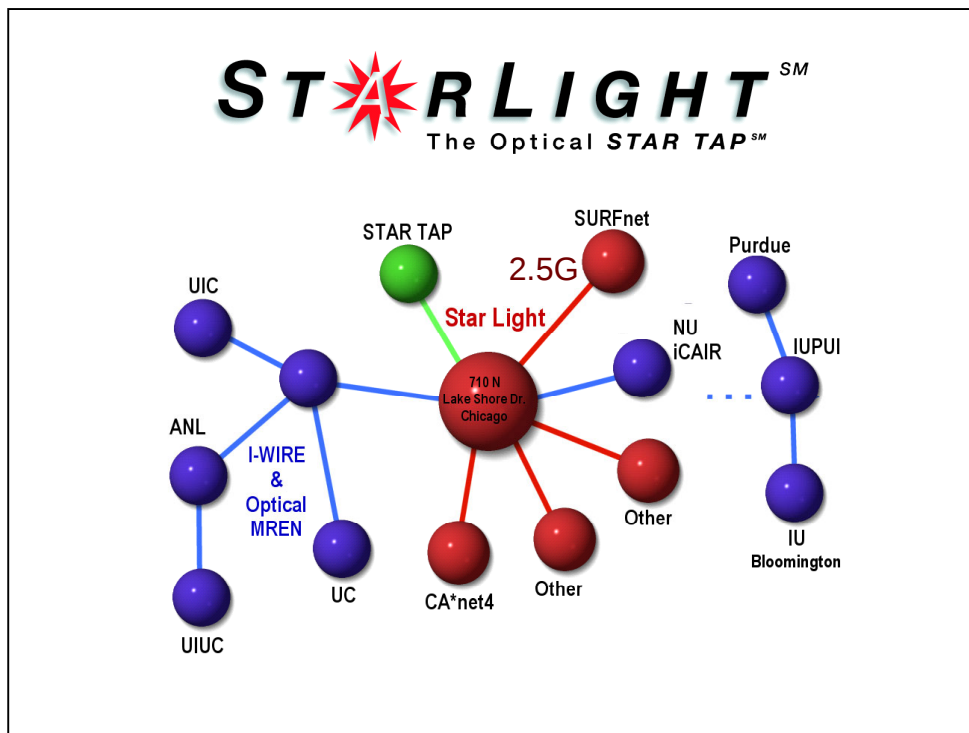


Figure 30

GRID CLUSTER

A 24-processor Linux cluster provides a grid presence for TRECC. Sponsored development initiatives include security and grid management applications.

Grid Cluster

- **Establish a Grid Presence at TRECC**
 - 24 Processor Linux Cluster
- **Participation in Grid Development**
 - Security Applications
 - Grid Management Tools





Figure 31

LCD DISPLAY CLUSTER

A 16-Processor Linux cluster drives commodity LCD panels to create a 5120 x 3840 (19.7M) pixel display for large-scale visualizations in a portable kiosk.

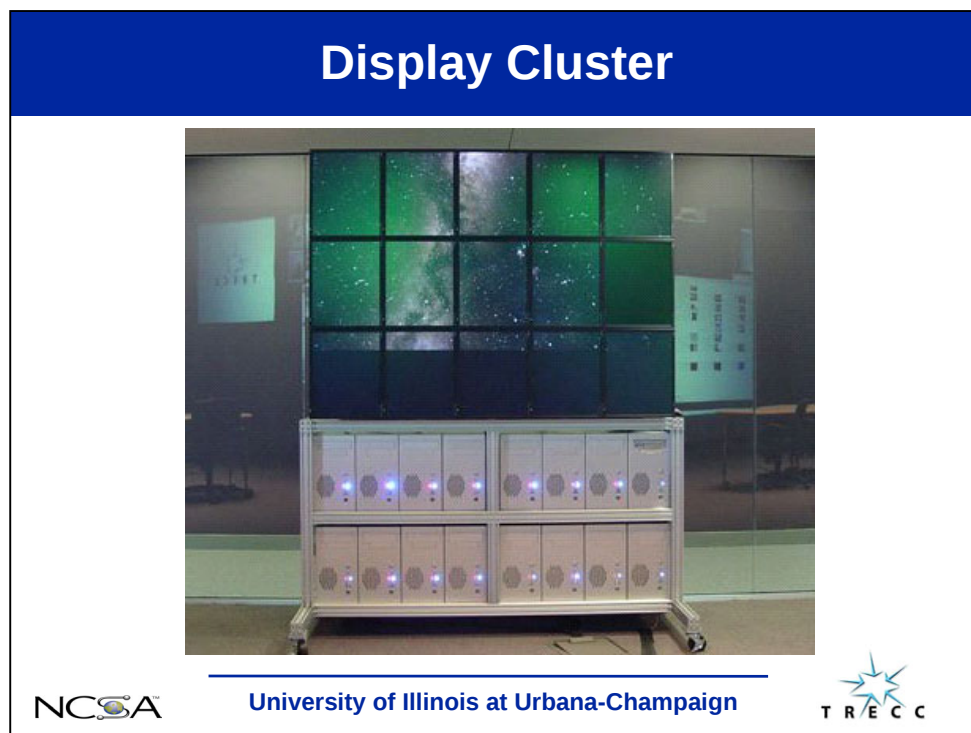


Figure 32

CONTINUUM

Continuum is a collaborative project between TRECC and the Electronic Visualization Laboratory (EVL) of the University of Illinois at Chicago to develop the hardware and software technology, and user-centered techniques for supporting intense collaborations in Amplified Collaboration Environments (ACE).

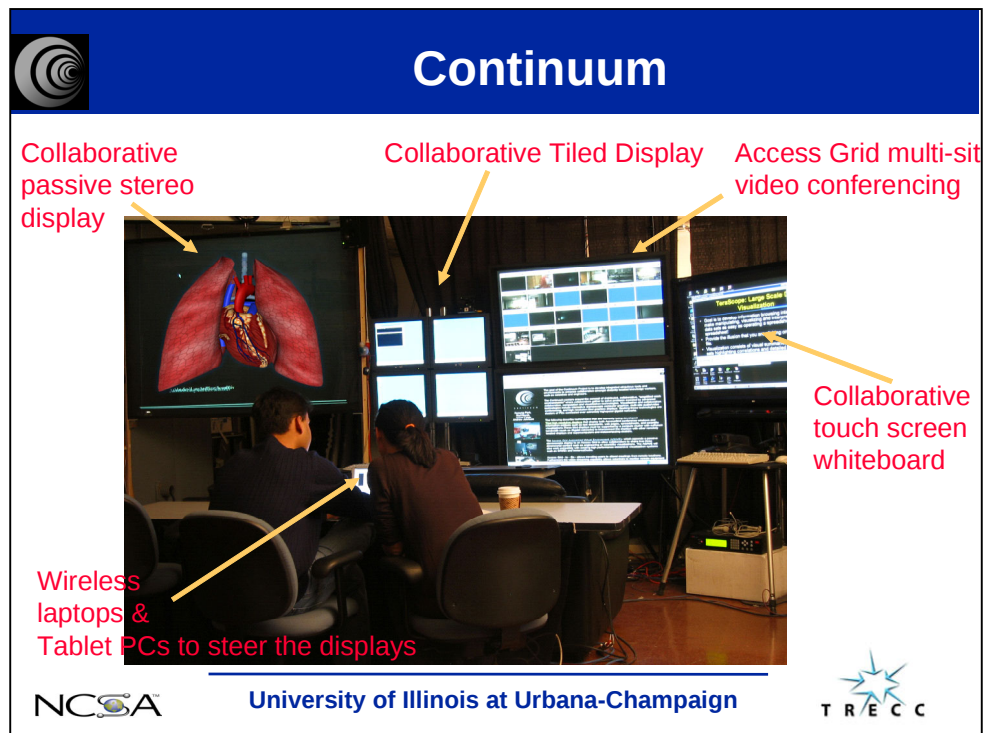
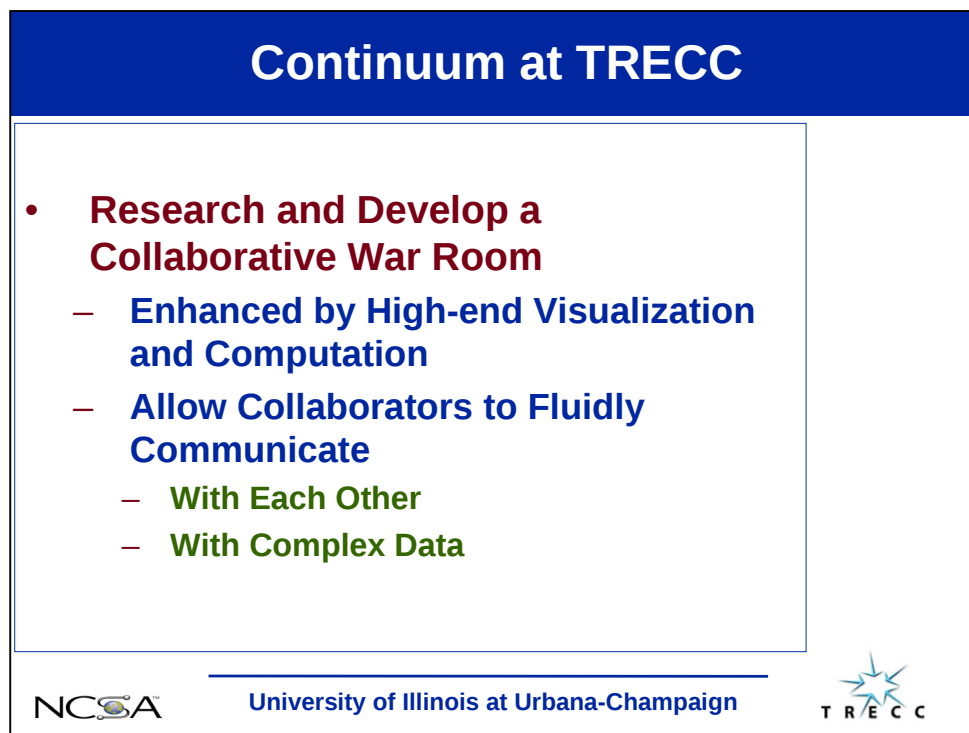


Figure 33

CONTINUUM GOALS

The goal of research in Amplified Collaboration Environments is to design future-generation collaboration spaces that take advantage of emerging advanced computing technologies, to allow collaborators that are geographically dispersed to work together as effectively as in traditional co-located project-rooms.

A presentation slide titled "Continuum at TRECC" with a blue header. The main content is a bulleted list in a white box. The list includes a primary goal in red and three sub-goals in blue and green. The footer contains logos for NCSA, the University of Illinois at Urbana-Champaign, and TRECC.

Continuum at TRECC

- **Research and Develop a Collaborative War Room**
 - **Enhanced by High-end Visualization and Computation**
 - **Allow Collaborators to Fluidly Communicate**
 - **With Each Other**
 - **With Complex Data**

NCSA University of Illinois at Urbana-Champaign TRECC

Figure 34

TELE-IMMERSION

The immersion module consists of an Access Grid Augmented Virtual Environment (AGAVE) passive stereo virtual reality display for visualizing three-dimensional data sets. The AGAVE is a low-cost VR system utilizing commodity conferencing projectors and a single PC equipped with modest dual-output graphics card.

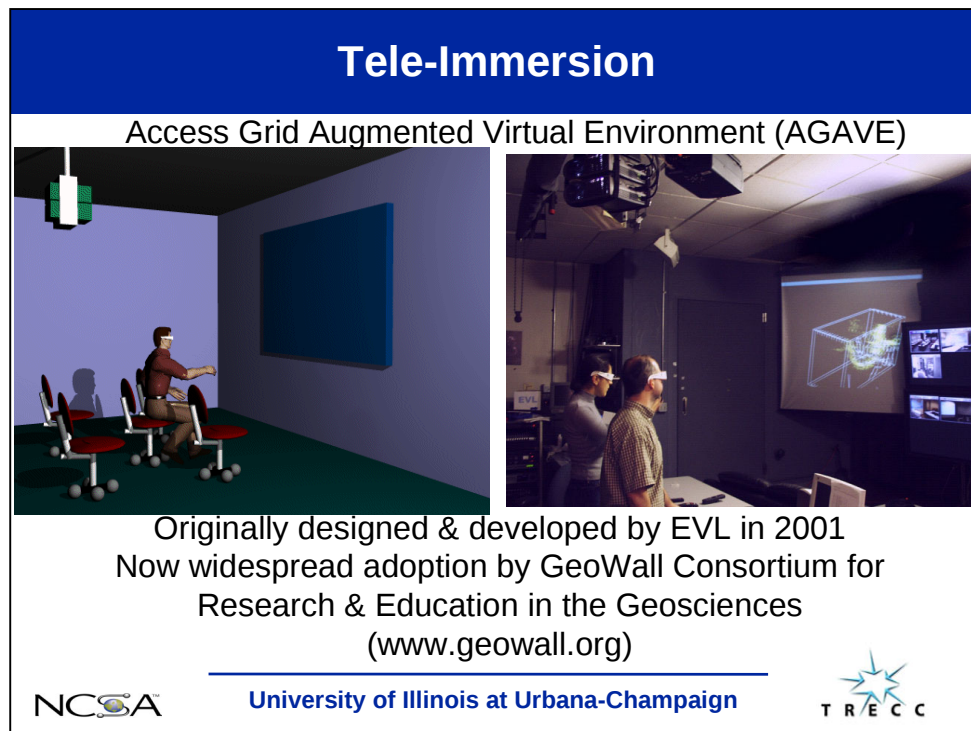


Figure 35

COLLABORATIVE ANNOTATION

An annotation module supports “whiteboarding” during a collaborative meeting. The technology employed is a plasma display enhanced with an touchscreen overlay to provide pen-based input.

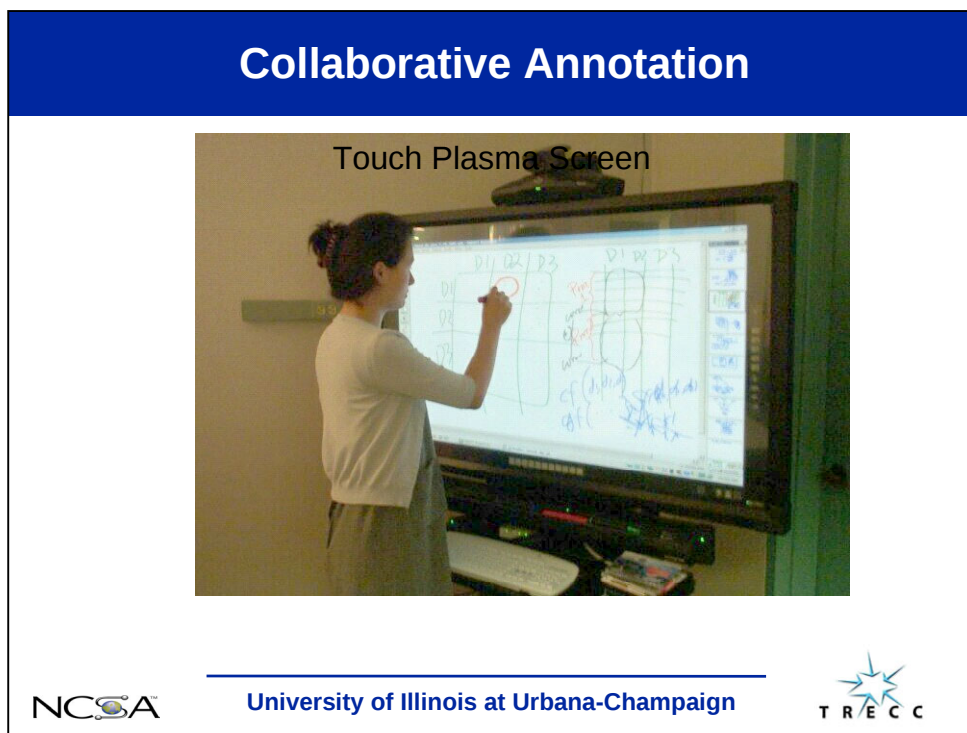


Figure 36

WIRELESS INTERACTION

The Continuum is controlled via a web interface using laptops and TabletPC's to steer data across a single seamless desktop. Work has begun to enable identification and authentication of users to tailor information to ones needs/interests.

2003 : Wireless Interaction



- Remote and Collaborative Steering of all Continuum Displays as One Screen.
- Push Wireless Devices Onto Tiled Displays.
- Walk in a Room & Start Sharing Information.

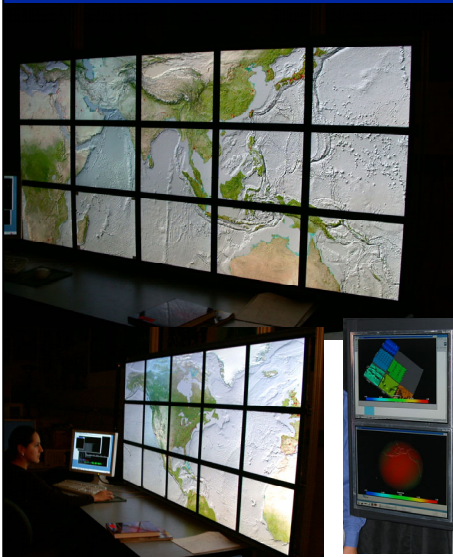


Figure 37

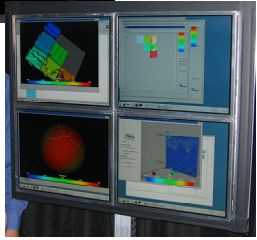
CONTENT DISTRIBUTION EXAMPLE

The content distribution module provides a scalable LCD tiled display for visualizing high resolution data sets.

JuxtaView : Extremely High Resolution Digital Montage Visualization for Tiled Displays



- Large Digital Montage Viewer for Tiled LCD displays
- View High Resolution Montages from Scripps, USGS
- USGS has aerial photos of 133 urban areas:
 - 5643 tiles each 5000x5000 pixel resolution ~ 375,600x375,600 pixels for each urban area (394GB per area.)
 - Total data ~ 51 TB



NCSA University of Illinois at Urbana-Champaign **TRECC**

Figure 38

TERAVISION

TeraVision is a way to remotely display moving graphics or high-definition video over gigabit networks. A basic system consists of a PC server with commodity hardware to grab high-resolution VGA or DVI inputs and a PC client to receive and display the streams.

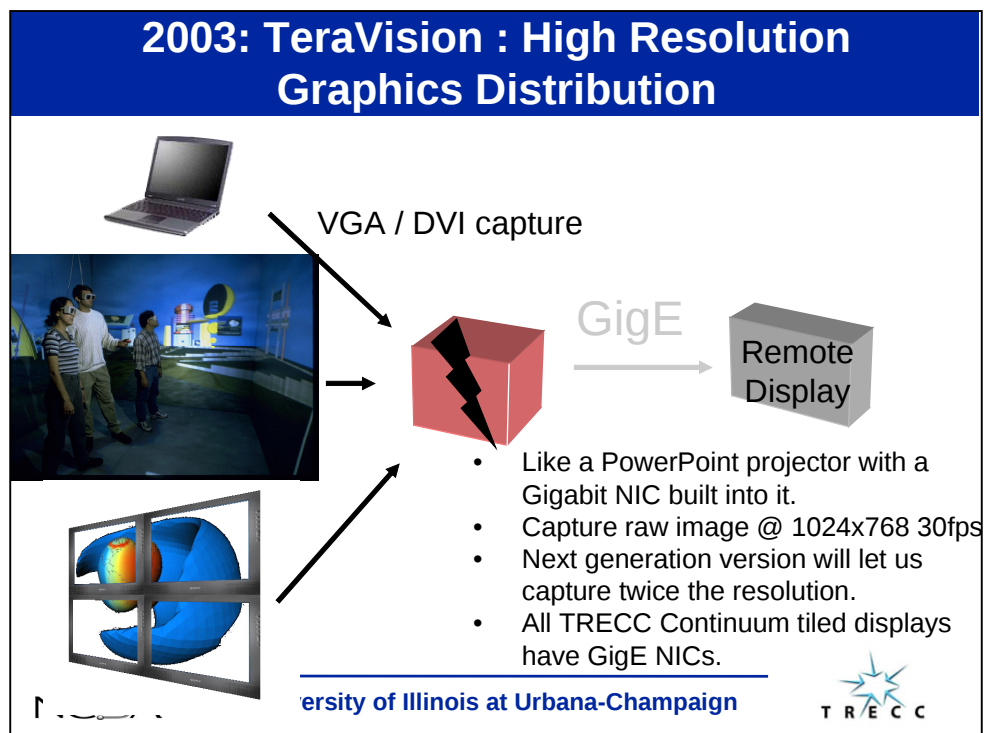


Figure 39

PARALLEL TERAVISION

Multiple TeraVision boxes can be used to stream component video streams of a tiled display. In a multicast network the streams can be efficiently shared in an N-way distribution.

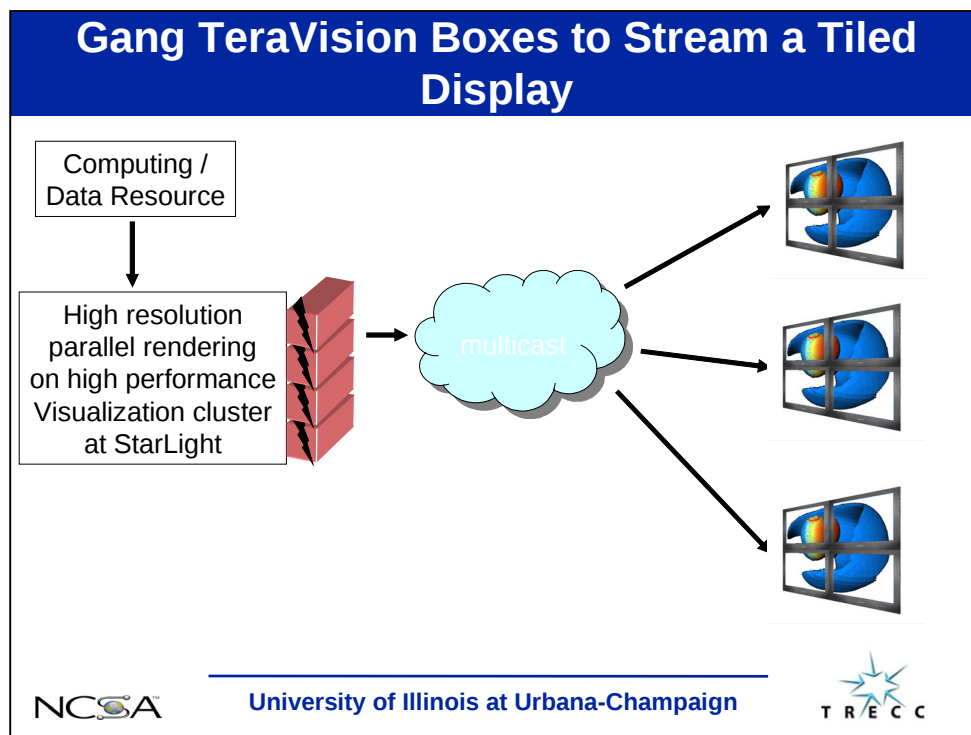


Figure 40

PARIS

Using PARIS (Personal Augmented Reality Immersive System), a projection-based “augmented” virtual reality display, a surgeon and medical modeler can view a three-dimensional model of a patient’s computed tomography (CT) data, and collaboratively review, sculpt and “virtually” build an implant using their hands.

PARIS is optimized to allow users to interact with the environment using a variety of tactile input devices. An authentic sensation of the implant sculpting process is achieved using SensAble Technologies’ PHANTOM force-feedback device. The PARIS display has excellent contrast and variable lighting that allows a user’s hands to be seen immersed in the imagery.



Figure 41

ACKNOWLEDGEMENTS

I would like to acknowledge the following contributors for information and slide materials.

1. Tate, G., TRECC General Briefing Package, Mar 2003.
2. Prudhomme, T., ONR Briefing, Nov 2000.
3. Wentling, T., A Knowledge Center Approach for Sharing Across Organizations, Mar 2003.
4. Meek, P., DPA Flight Center, Mar 2003.
5. Thot-Thompson, J., Cyberinfrastructure Research for Homeland Security, Feb 2003.
6. Marcusiu, D., TRECC Quarterly Review, Jul 2002.
7. Leigh, J., The Continuum Project: Research in Amplified Collaboration Environments, Mar 2003.

Acknowledgements	
• Gail Tate – TRECC • Tom Prudhomme – NCSA • Tim Wentling – NCSA • Pamela Meek – DPA • Janet Thot-Thompson – ACCESS DC • Doru Marcusiu – NCSA • Jason Leigh – EVL	
	
	

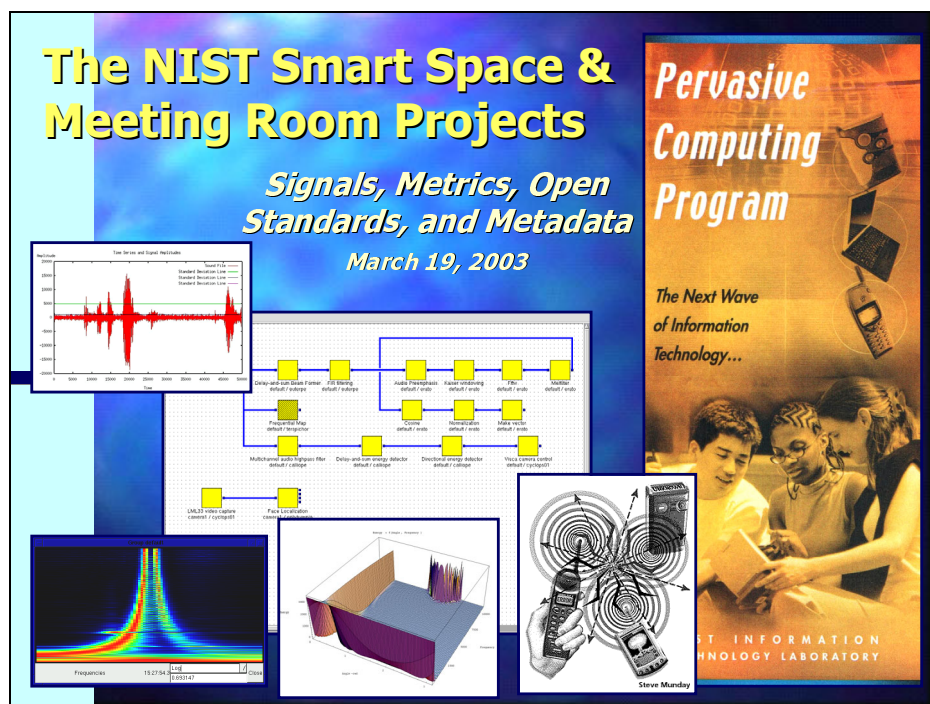
Figure 42

THE NIST SMART SPACE & MEETING ROOM PROJECTS

Vincent Stanford
National Institute of Standards and Technology
Gaithersburg, MD

THE NIST SMART SPACE & MEETING ROOM PROJECTS

Pervasive computing devices, sensors, and networks, provide infrastructure for context-aware smart meeting rooms that sense ongoing human activities and respond to them. These technologies require advances in areas including networking, distributed computing, sensor data acquisition, signal processing, speech recognition, human identification, and natural language processing. Open interoperability and metrology standards for the sensor and recognition technologies can aid R&D programs in making these advances. To address this need the NIST Smart Space and Meeting Room projects are developing tools for data formats, transport, distributed processing, and metadata. We are using them to create annotated multi modal research corpora and measurement algorithms for smart meeting rooms, which we are making available to the research and development community.



Slide 1

NEW INTERFACES – VISIONS AND CHALLENGES

Visionary system concepts presented at this conference, like the MIT Oxygen project, DARPA High Productivity Computing Systems program, the IBM Pervasive Computing initiative, and NASA cognitive systems concepts offer approaches to high productivity in collaborative aerospace engineering design and development organizations. Most of these include sensor based statistical recognition systems for speech, speakers, faces, gestures, and even emotional states of users. These can be combined into a perceptive interface that has a sense, recognize, understand, and respond cycle of operation. The understanding component differentiates a perceptive interface from a more traditional stimulus-response perceptual interface. Prerequisite sensor systems, such as advanced microphone arrays, to provide adequate signal quality for these recognition algorithms are now under development. However, the recognition systems needed still have significant error rates and will have to be incrementally improved and made more robust with respect to environmental conditions.

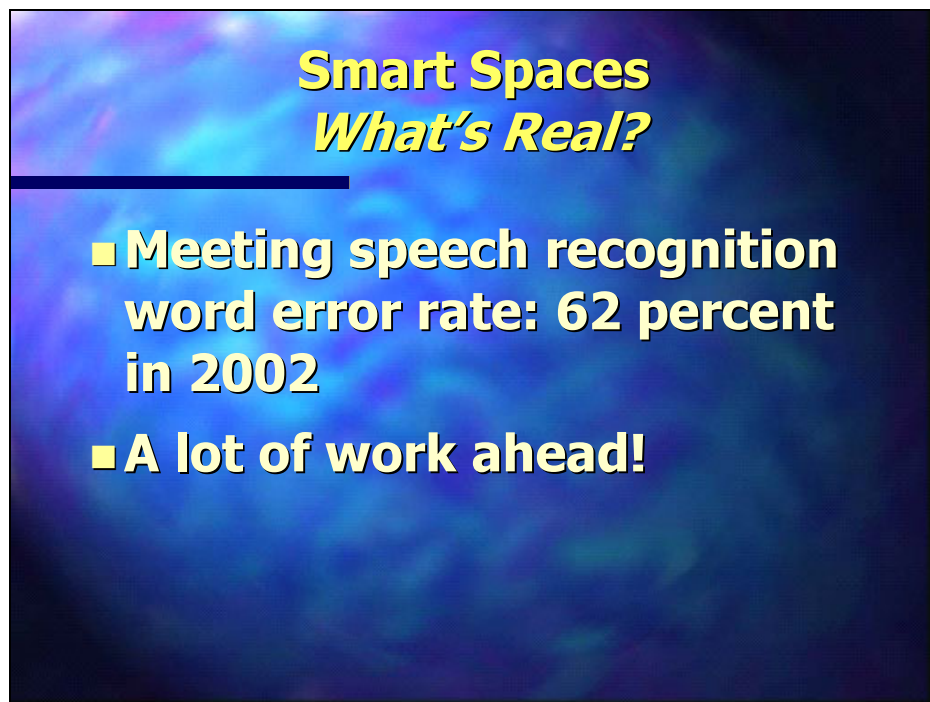
New Interfaces *Visions and Challenges*

- Visionary system concepts, like oxygen, HPCS, Cognitive, and Pervasive systems are essential road maps; or under construction
- But real challenges remain to developing systems that:
 - ***Sense*** user signals like speech, gesture, and physiological measurement
 - ***Recognize*** words, speakers, gestural referents
 - ***Understand*** context, and user intent
 - ***Respond*** with information retrieval, computation, and rendering

Slide 2

SMART SPACES – WHAT'S REAL?

The vision of systems that can respond to spontaneous speech that is inferred from existing commercial large vocabulary dictation systems is optimistic. While they do perform adequately for highly codified technical speech, say radiology dictation by a physician, they do not yet provide good recognition performance for spontaneous speech. This fall, NIST conducted an evaluation of a state-of-the-art large vocabulary speech recognition system, and found that it missed two words out of three for spontaneous person to person speech in meetings. This area still requires additional investigation and development, with others such as face recognition, under actual field conditions, having similar performance problems.



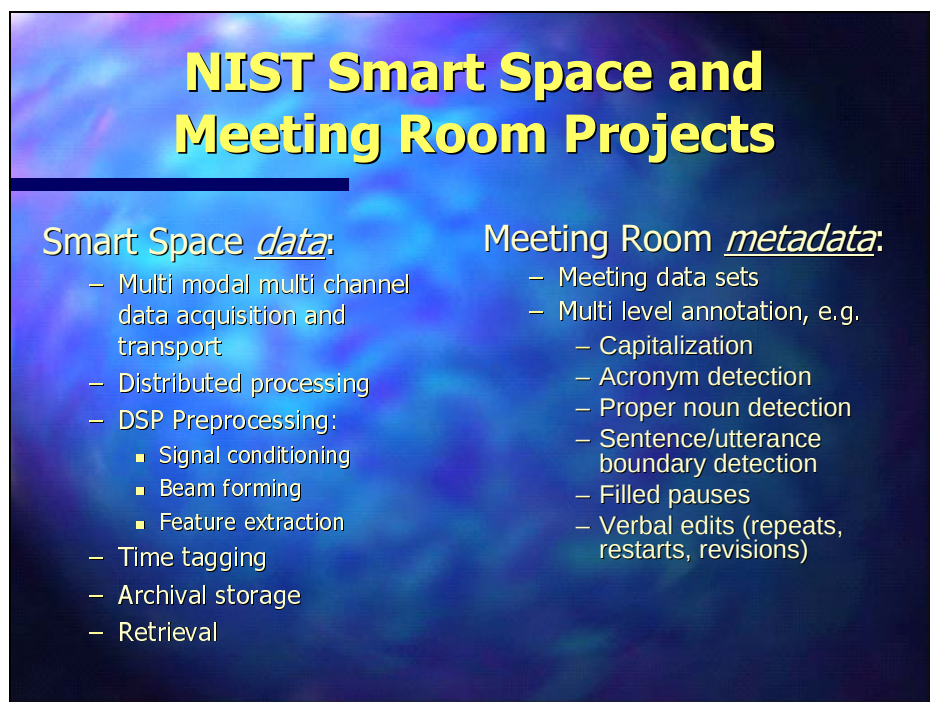
Smart Spaces
What's Real?

- Meeting speech recognition word error rate: 62 percent in 2002
- A lot of work ahead!

Slide 3

NIST SMART SPACE AND MEETING ROOM PROJECTS

We address issues of data acquisition using the NIST Smart Data Flow System system, which is a set of tools that allow components from various developers to interoperate in an environment containing flows from many sensors, and offering a reference implementation for laboratory use. The NIST Meeting Room has over two hundred microphones, five cameras, a smart whiteboard, and will soon have a locator system for the meeting attendees. In the aggregate, these generate over a gigabyte per minute of sensor data, which are time tagged to millisecond resolution and stored for research uses. We address broad issues of metadata, or annotations, with semantic descriptions using the Architecture and Tools for Linguistic Analysis Systems (ATLAS). One of the major design features of ATLAS is standardization of metadata derived directly from the sensor data streams, and subsequent higher-level annotations of meeting context, which may allow indexing, transcription, and possibly even summarization of meetings. Some significant meeting metadata under investigation include spoken words, speaker identity, sentence-like units, disfluencies, speaker locations, and time tags. From these low-level metadata, smart spaces will have to make higher-level inferences about tasks the users are undertaking to become context-aware.



NIST Smart Space and Meeting Room Projects

Smart Space <i>data</i>:	Meeting Room <i>metadata</i>:
– Multi modal multi channel data acquisition and transport – Distributed processing – DSP Preprocessing: ▪ Signal conditioning ▪ Beam forming ▪ Feature extraction – Time tagging – Archival storage – Retrieval	– Meeting data sets – Multi level annotation, e.g. – Capitalization – Acronym detection – Proper noun detection – Sentence/utterance boundary detection – Filled pauses – Verbal edits (repeats, restarts, revisions)

Slide 4

ON THE WAY TO UNDERSTANDING...METADATA RICH TRANSCRIPT

Perceptual interfaces will allow smart meeting rooms to act as meeting secretary by taking meeting minutes from the chairperson in response to commands. The NIST Rich Transcription Evaluation series, which began in 2002, seeks to support the development of these technologies. A raw machine generated transcript, XML metadata enrichments, and human-readable form are shown here to clarify the importance of such annotation capabilities. Future cognitive and perceptive interfaces will have to extract meaning from word streams generated by speech recognizers. An initial challenge will be creation of rich transcripts that are easily readable by humans. This will require automated processing of speaker terms, named entity tracking, and later, even topic identification. We believe that a long term program providing standard reference materials, metrics, and algorithm evaluation will be needed to enable the creation of the usable and facile multi modal interfaces of the future.

**On the Way to Understanding...
Metadata Rich Transcript**

ASR Symbol Stream	XML Annotations	Readable Transcript
tonight this thursday big pressure on the clinton administration to do something about the latest killing in yugoslavia airline passengers and outrageous behavior at thirty thousand feet what can an airline do and now that el nino is virtually gone there is la nina to worry about.	<pre><speaker name="Peter Jennings"> <sent type=decl> tonight this <proper_noun> thursday </proper_noun> <phrase- break> big pressure on the <proper_noun>clinton </proper_noun> administration to do something about the latest killing in <proper_noun> yugoslavia </proper_noun></sent> <sent type=decl> airline passengers and outrageous behavior at <numex val=30,000> thirty thousand </numex> feet </sent><sent type=inter> what can an airline do</sent> <sent type=decl> and now that <proper_noun> el nino</proper_noun> ...</pre>	<u>Peter Jennings:</u> Tonight, this Thursday, big pressure on the Clinton administration to do something about the latest killing in the Yugoslavia. Airline passengers and outrageous behavior at 30,000 feet. What can an airline do? And now that El Nino is virtually gone, there is La Nina to worry about.

Slide 5

SMART SPACES – WHAT'S REAL?

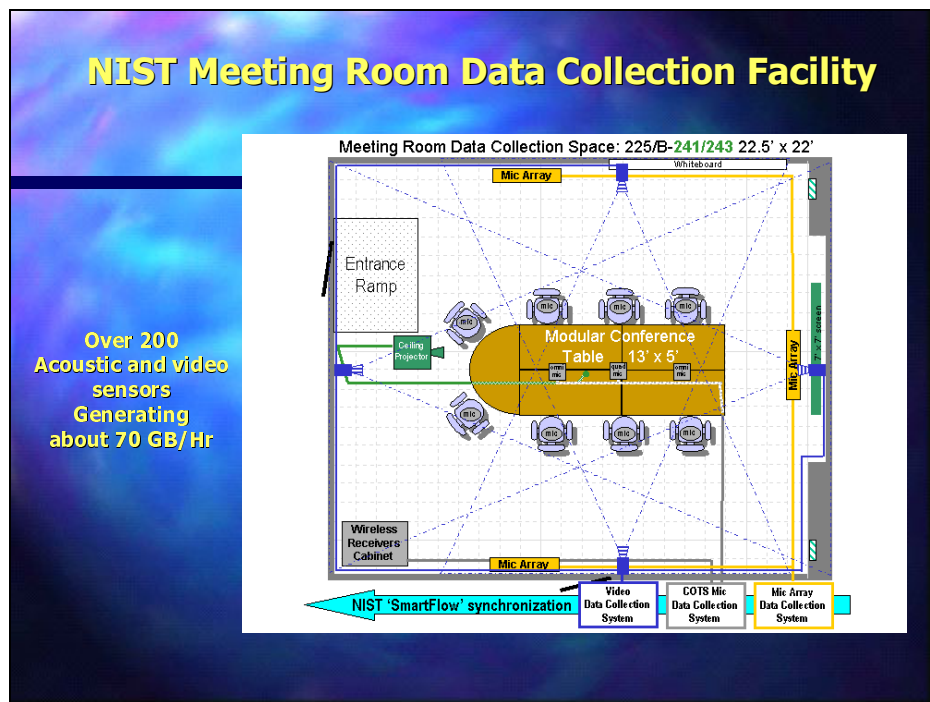
With current technology can acquire speech at a distance using a phased microphone array, and perform speaker dependent speech recognition. This currently requires a skilled and cooperative speaker, and discourse in the domain of the language model. Recent work shows that it is possible to apply a GMM based speaker verification algorithm to the cepstral coefficients used for speech recognition to also identify the speaker. This will allow near real time speech recognition for a privileged user in a meeting space. The recognized speech can be used to transcribe meeting minutes, or to voice commands to the system. This may allow useful capabilities to be deployed long before spontaneous speech and natural language understanding emerge.

Smart Spaces – What's Real?

- Speech recognition possible using microphone arrays *for skilled speakers*
- Speech segmentation and speaker verification possible
- A selected skilled user can be transcribed in a cooperative group
- Transcribed speech can be parsed for basic commands

NIST MEETING ROOM DATA COLLECTION FACILITY

A schematic plan view of the layout and sensor arrangements in the NIST Meeting Room Laboratory is shown above. It is a sensor rich environment which provides many views of the meetings using twenty-four random placement microphones, three linear microphone arrays, five camera views, and an electronic white board. We are currently developing enhancements to this facility. One being our Mk-III microphone array, which offers improved signal to noise ratio, and onboard conversion of the data to UDP/IP packets for direct transport across fast Ethernet. Also, an additional sensors for meeting participant locations using smart badges is being added. The NIST Meeting Room project currently uses this facility to record meetings of small groups, and offers the data to research and development communities.



Slide 7

MULTI MODAL MEETING RECORDING

We have recorded meeting room data for industry and academic research and development groups. This consists of twenty hours of meeting data, at more than seventy gigabytes per hour. These meetings had various subjects including focus groups, game playing, expert interviews, and planning meetings. They varied in length from fifteen minutes to one hour, and had from three to eight participants. This data will be made available through the Linguistic Data Consortium.

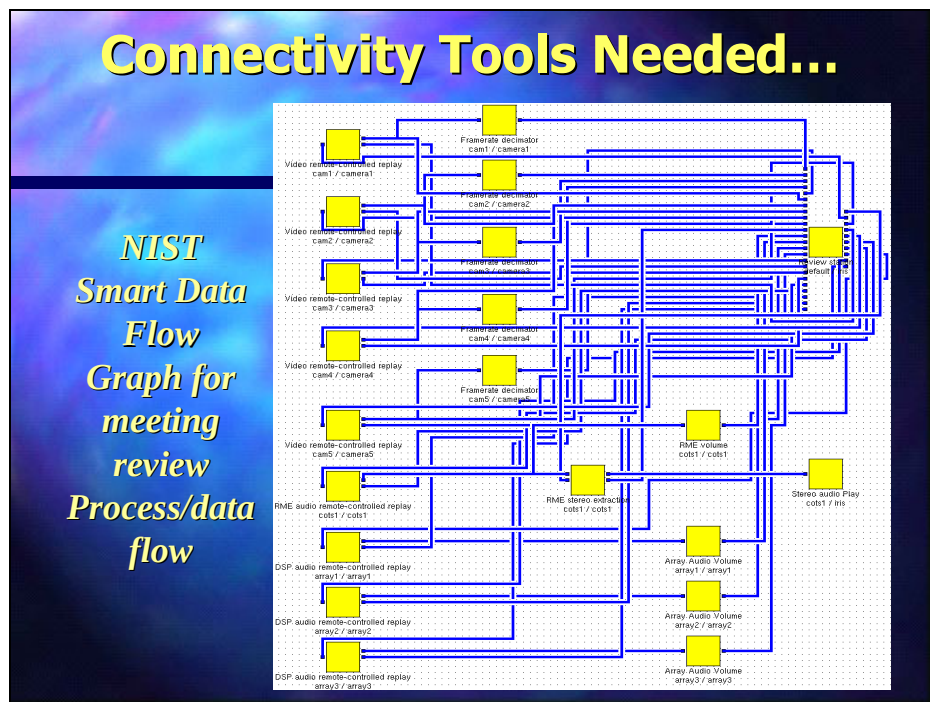
Multi modal Meeting Recording

- Collect and review recordings
- Open system-based, interfaces with Smart Data Flow live or from archived data
- User selects video views and audio channels
- User controls camera view/movement

Slide 8

CONNECTIVITY TOOLS NEEDED...

The NIST Smart Data Flow System was developed in response to the need to provide connectivity to the large number of sensors and devices that will be needed to construct smart meeting rooms, and perceptual interfaces. An operational flow graph for data review in the NIST Meeting Room is shown above. The Smart Data Flow System generates the connections and transports the data among the clients represented by the blocks. Dragging and dropping the components and naming of the flows can be used to reconfigure the application flow graphs. The system consists of a defined middleware API for real-time data transport, and a connection server for sensor data streams.



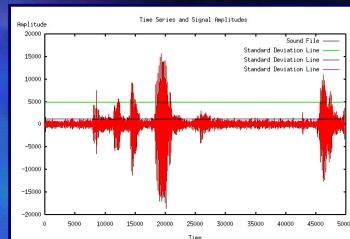
Slide 9

METRICS RESEARCH NEEDED

Significant research is required in order to construct measurement protocols for mixed initiative systems that allow multiple actions, choices, and responses. Experience in the speech recognition community showed that well drawn measurement programs are very important to ongoing technical improvements in a new technologies. Reference data sets will be needed for the several recognition tasks required in Smart Spaces, and Smart Spaces can provide a test bed for integrated functionality of the technologies. Measurement tasks for multiple cascaded technologies will have to be designed.

Metrics Research Needed:

- **Training/test reference data:**
 - Acoustic source location
 - Speech recognition
 - Biometric speaker ID
 - Face recognition
- **End-to-end metrics for perceptual interfaces**
 - Sensing – SNR, timing, etc
 - Recognition – Who? What's going on? Where? What were they looking at?
 - Response – How can the system help?
- **Information retrieval and access tasks**
- **Distributed interfaces using pervasive devices**



MEETING ROOM DATA COLLECTION LABORATORY

The NIST Smart Space employs a combination of software and hardware to provide a test bed for sensor based interface components. It provides data acquisition, archiving, time tagging, and data transport for the recognition components that will make up the multi modal collaborative interfaces envisioned in this conference. Some examples of possible technologies are listed above.

Meeting Room Data Collection Laboratory

- **Phased microphone arrays**
- **Computer-controlled video cameras**
- **Biometric sensor fusion using commercial components:**
 - **Acoustic speaker identification**
 - **Facial image classification**
- **Speaker dependent speech recognition**
- **Data flow test-bed for integration of commercial products**

Slide 11

NIST SMART DATA FLOW MIDDLEWARE

The NIST Smart Data Flow System is a proposed reference implementation of data transport and format standards for sensor intensive interfaces. Our prototype has been tested in the NIST science programs. Such systems must be distributed across numerous nodes, provide standardized data transport mechanisms and data types, but allow for more to be defined. They must also abstract connectivity mechanisms and support device/service discovery, be fault tolerant, and allow mobile nodes to come and go from the environment.

[illegible]

MULTI MODAL PROCESSING

Some examples of visual interface processing might include face localization using skin color detection, face normalization using reference points like eye locations, and gesture recognition. Acoustic interface components, based on phased array processing, include source location, speech recognition, speaker identification, and sensor fusion with visual data.

Multi Modal Processing

Mark-I NIST Microphone array

Time delay matrix

X(0) → ADC
Y(0) → ADC
X(1) → ADC

Time

Distance

Beam

Different look directions
each arrow represents a steering
direction for a different look angle

coords [Quarter Grayscale] Close

brow [Quarter Grayscale] Close

feat [Quarter Grayscale] Close

Slide 13

USABILITY FEATURES NIST SMART DATA FLOW SYSTEM

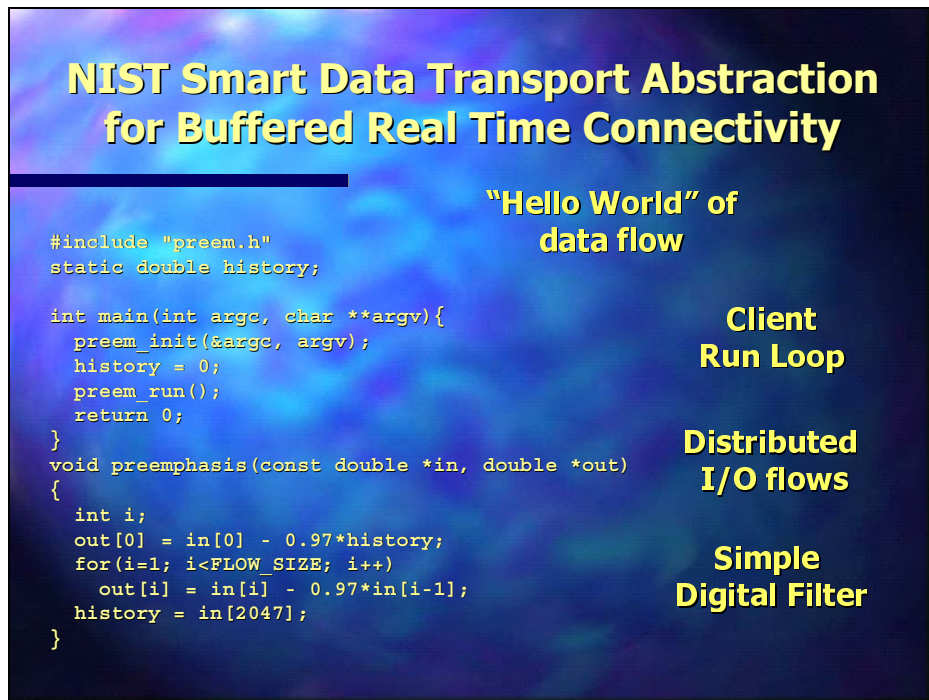
We have worked with several laboratories to understand problems encountered when building sensor intensive smart spaces. We found that our users are interested in a more streamlined capability. We are forming an industry working group to work on further definition of standards that will facilitate integration and testing of the many technology components necessary to implement the advanced mixed initiative systems envisioned for the future.

Usability Features NIST Smart Data Flow System

- **Initial version was difficult to deploy and use – New version under development:**
 - Visual flow graphs
 - Code generator
 - Simplified API
 - Device, user, and service discovery
 - Fault tolerant

NIST SMART DATA TRANSPORT ABSTRACTION FOR BUFFERED REAL TIME CONNECTIVITY

The NIST Smart Data Flow system provides an abstraction for connectivity and data buffering to facilitate the construction of the needed multi process, distributed, systems. The pointers to data flows “in” and “out” can reside on remote systems with the connectivity being defined at the graphical level. This allows data flow component libraries that integrate various real time signal processing and recognition capabilities to be defined and used in a variety of flow graph contexts. Substantial code is generated by the data flow middleware to support this simplified application structure.



**NIST Smart Data Transport Abstraction
for Buffered Real Time Connectivity**

```
#include "preem.h"
static double history;

int main(int argc, char **argv){
    preem_init(&argc, argv);
    history = 0;
    preem_run();
    return 0;
}

void preemphasis(const double *in, double *out)
{
    int i;
    out[0] = in[0] - 0.97*history;
    for(i=1; i<FLOW_SIZE; i++)
        out[i] = in[i] - 0.97*in[i-1];
    history = in[2047];
}
```

**"Hello World" of
data flow**

**Client
Run Loop**

**Distributed
I/O flows**

**Simple
Digital Filter**

Slide 15

NIST MARK-III MICROPHONE ARRAY

NIST has been involved in developing spoken language corpora to support the training and testing of large vocabulary speech recognition systems. These data sets progressed from a one-thousand word structured task, to five-thousand word vocabulary readings, to twenty-thousand word readings, to recorded broadcast news programs, and most recently to spontaneous speech in small group meetings. The meeting room laboratory offers multiple views of speech at progressively greater distances with close-talk microphones at one or two inches from the lips, lapel microphones, table top microphones, and wall mounted microphone arrays. We have distributed this technology to interested research and development laboratories, such as the Georgia Institute of Technology Aware Home project, shown above in the New York Times on April 5, 2001. The array technology shown is of the Mark-II series, which was our first digital phased array.



Slide 16

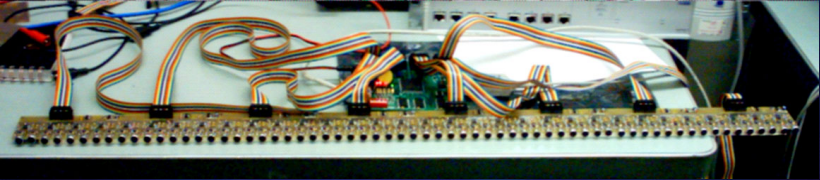
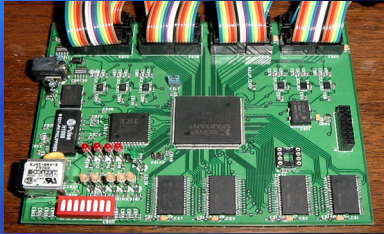
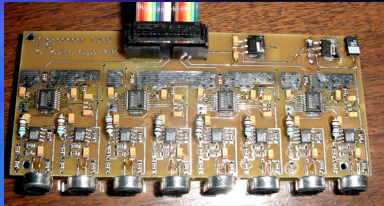
THE MARK-III MICROPHONE ARRAY

An operational prototype of the Mark-III series microphone array is now completed, and has important advantages in terms of manufacturability, signal quality, price, and deployability in research environments. We are making the construction data available to interested research and development laboratories. This mark uses twenty-four bit analog to digital conversion, an Ethernet based data interface, and a local field programmable gate array to read the ADCs, and create TCP/UDP frames that are sent them to other Smart Data Flow Nodes. This hardware architecture consists of a motherboard which reads digital interfaces from each of eight daughter cards with eight microphones. This mark has improved signal to noise characteristics due to the short analog signal runs, and improved ADC resolution and noise floor.

The Mark-III Microphone Array

Integrated, Easy to Replicate

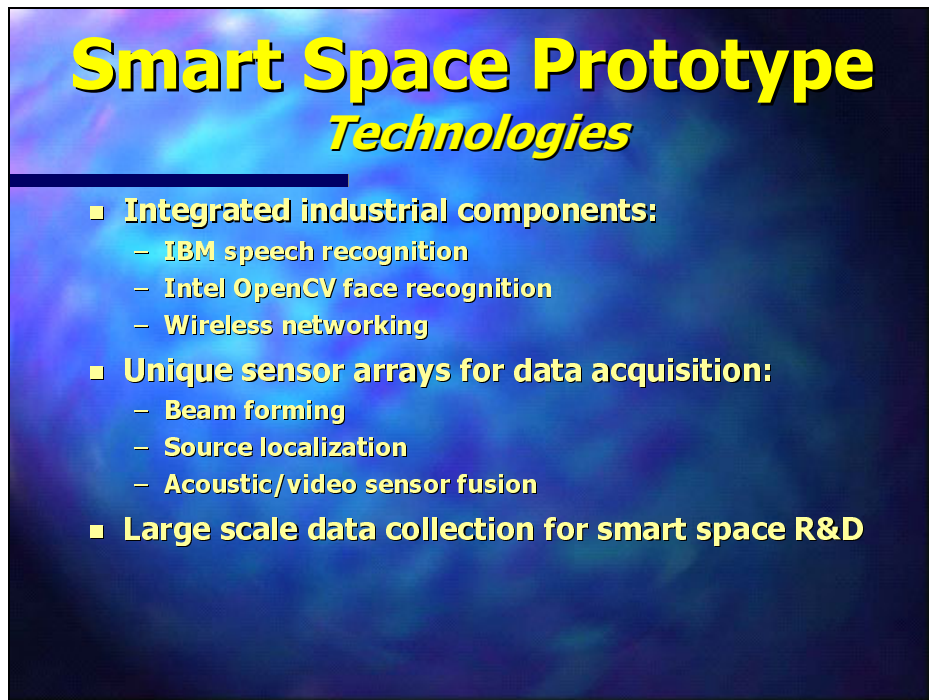
- VLSI, FPGA, VHDL, Preamps, ADCs
- 64 microphones, 24bits at 48kHz
- 2Mbyte local data buffer
- IP negotiation with a BOOTP server
- Data transmission via Fast Ethernet
- Responds to simple requests from Smart Data Flow: array ID, receiving node, array active indicator, etc.



Slide 17

SMART SPACE PROTOTYPE TECHNOLOGIES

Our Meeting Room apparatus consists of microphone arrays acquiring sixty-four channel audio input flows and offering them for subscription. A beamformer subscribes to these, reduces them to a single channels, offering audio flows. Many sensor and recognition technologies are under development in industry, so interoperability and integration issues are crucial to new generation smart environments. The NIST Smart Data Flow System is being used to integrate technologies including: speech recognition, speaker identification, face, localization and recognition, channel normalization, video and acoustic displays, and wireless PDAs. Standardized formats are offered for multimedia data streams, archiving, retrieval, and review tools. Hand crafting the needed inter-process communication was found to be very labor intensive and brittle with respect to changing requirements for new sensors and configuration changes to accommodate equipment faults. The NIST Smart Data Flow System toolkit has components for graphical configuration of flows, allocation of the graph nodes to distributed systems, and connection by TCP/IP/UDP. Data transport code is provided by the Smart Data Flow System libraries. We hope to make this the basis of a standards working group and collect additional requirements from industry and work cooperatively to develop reference implementations for smart meeting and multi modal interfaces.

A presentation slide with a dark blue background and a lighter blue abstract pattern. The title "Smart Space Prototype Technologies" is at the top in large, bold, yellow font. Below the title, there are three main bullet points, each preceded by a yellow square. The first bullet point is "Integrated industrial components:" followed by three sub-bullets: "IBM speech recognition", "Intel OpenCV face recognition", and "Wireless networking". The second bullet point is "Unique sensor arrays for data acquisition:" followed by three sub-bullets: "Beam forming", "Source localization", and "Acoustic/video sensor fusion". The third bullet point is "Large scale data collection for smart space R&D".

Smart Space Prototype Technologies

- Integrated industrial components:
 - IBM speech recognition
 - Intel OpenCV face recognition
 - Wireless networking
- Unique sensor arrays for data acquisition:
 - Beam forming
 - Source localization
 - Acoustic/video sensor fusion
- Large scale data collection for smart space R&D

Slide 18

SENSORS WILL ALLOW PERSONAL INTERFACES

The multi modal, recognition based, interfaces of the near future will allow personalized interfaces to respond to selected individuals, maintain user profiles and session histories to provide some degree of context awareness. This will enable the interactive, mixed initiative, project design and management environments envisioned in this conference.

Sensors Will Allow Personal Interfaces

- Current interfaces are *nominal* – who presses the buttons does not matter
- Sensor based interfaces can be *personalized*
 - Recognize – *who said what, gestures etc.*
 - Understand – *what did it mean in context*
- “Computer, bring up my appointment calendar.”

VISION OF THE POSSIBLE: A MEETING ROOM THAT...

We are proposing an integration challenge to the providers of the statistical recognition software and other relevant components. It is designed to be at the edge of the current state of the art. It would consist of a meeting room using microphone array technology that is sensitive to a meeting chairman in particular over the other meeting participants. This will require integration of several technology components in cascade and parallel to provide the necessary signal acquisition, conditioning, preprocessing, recognition, and responses. Such a multi modal system could also serve as the foundation of accessible computing, with standards for identifying user preferences and needs, and protocols to communicate them to host smart environments.

Vision of the Possible: *A Meeting Room That...*

- **Takes the minutes** from the moderator
- **Responds to commands**, depending on who spoke, what they were looking at, or pointing to
- **Accesses information** by voice query
- **Provides security** based on participant identity

Slide 20

ACCESSIBILITY PROTOTYPE: HANDS FREE SERVICES

A near term, and humane, use of the technology involved in the integration challenge would be to provide accessible computing. For example if a user wants, or needs to operate the computing, meeting, and presentation environment hands free. This will require user and service discovery, using dynamic networking and appropriate security safeguards. An initial prototype includes a PDA with 802.11 wireless networking that negotiates for services, communicates user preferences, and uploads personal profiles for speech, speaker, and possibly other recognition training data.

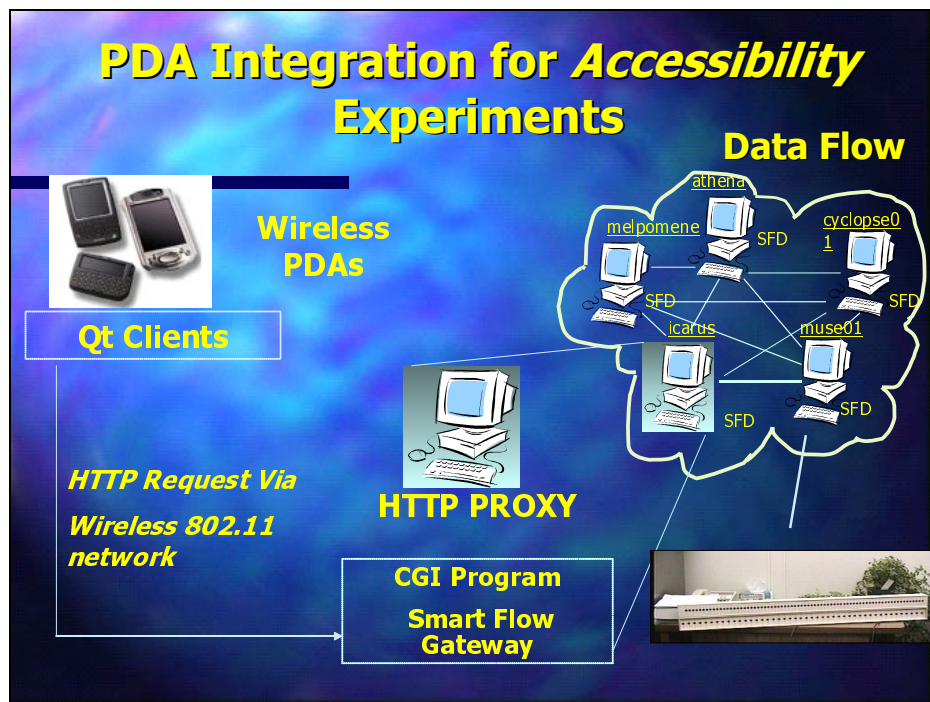
Accessibility Prototype: *Hands Free Services*

- User device discovery
- Hands free preference negotiation
- Service discovery:
 - Microphone array
 - Speaker ID
 - Speaker dependent speech recognition
- Upload biometric profiles for recognition
- Distributed data acquisition and processing

Slide 21

PDA INTEGRATION FOR ACCESSIBILITY EXPERIMENTS

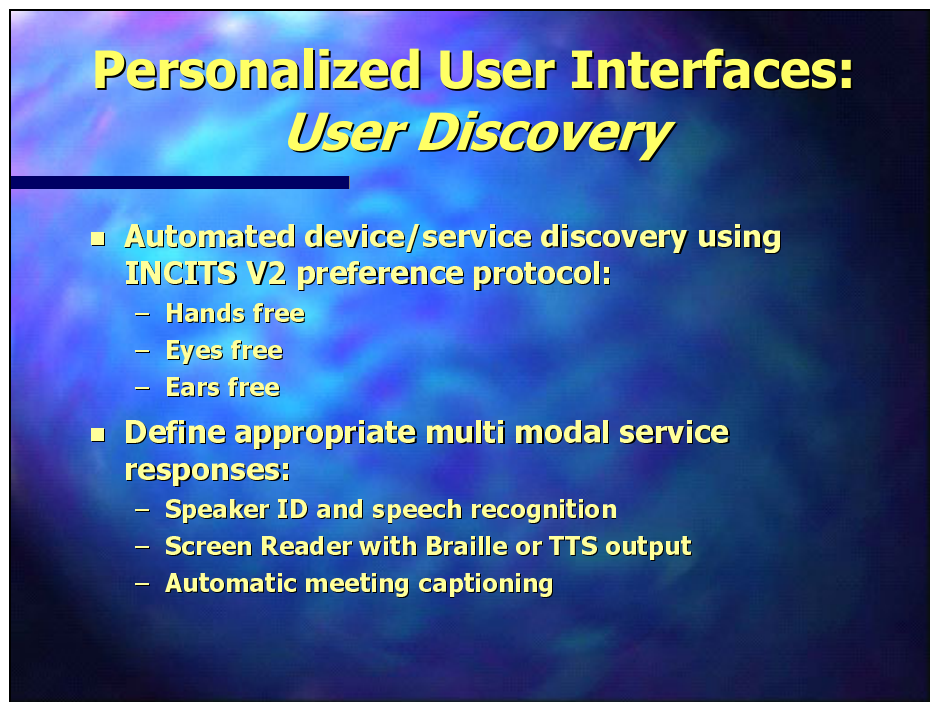
A prototype device discovery protocol includes wireless PDAs with 802.11 networking, and uses services including, DHCP, HTTP, and CGI, as well as INCITS-V2 protocols to communicate user preferences, and to initiate the required multi modal interface services and applications.



Slide 22

PERSONALIZED USER INTERFACES: USER DISCOVERY

The device/service discovery protocols will allow user preferences like hands-free, eyes-free, and ears-free operation to the multi modal service environment provided by smart spaces. The NIST Smart Data Flow System can be used to integrate real time services that support speech recognition, activate screen readers, or closed captioning as specified in user preferences. Other service graphs can be defined and used to respond to additional preferences as the technologies emerge.



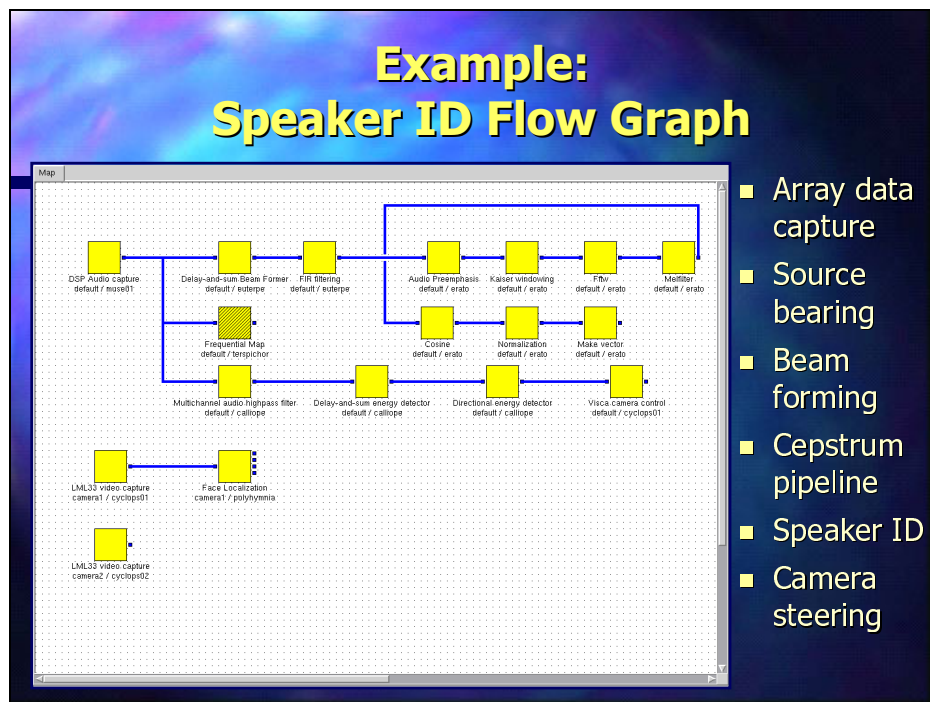
Personalized User Interfaces:
User Discovery

- **Automated device/service discovery using INCITS V2 preference protocol:**
 - Hands free
 - Eyes free
 - Ears free
- **Define appropriate multi modal service responses:**
 - Speaker ID and speech recognition
 - Screen Reader with Braille or TTS output
 - Automatic meeting captioning

Slide 23

EXAMPLE: SPEAKER ID FLOW GRAPH

An example of a distributed flow graph that could provide some components needed is shown above. It captures data from a NIST microphone array, and sends it to a source bearing estimator, a beam former, and a computer controlled camera that can point to a speaker. This graph is currently operational in the NIST Smart Space Laboratory on an experimental basis. Additional components could be integrated that make use of the video, and audio services available in the existing smart space framework.



Slide 24

WHAT CAN NIST DO FOR THIS COMMUNITY?

NIST is interested in aiding U.S. industry through the use of measurements and standards. Our Smart Space and Meeting Room projects offer metrology, reference data, and proto-standards for data transport and formats. We can also participate in standards working groups and publish non-regulatory standards in aid of industry groups to promote interoperability.

What Can NIST Do for this Community?

- Neutral entity chartered to enhance industry productivity through measurements and standards
- Can publish non-regulatory standards embodying community agreements
- Seeking industry/academic partners
- Advanced Metrology, physical and information sciences
- Cooperatively produce standard reference data sets
- Make measurement algorithms and protocols publicly available

Slide 25

MEASUREMENTS AND STANDARDS WILL BE KEY...

To summarize: we believe that the development of the advanced cognitive interfaces discussed at this workshop can be facilitated by standardization and performance metrics. We would like to discuss the matter with interested parties in the research and development communities in industry, academic, and government laboratories.



Measurements and Standards Will be Key...

- **Performance metrics**
- **Standardized:**
 - Data formats
 - Transport mechanisms
 - Distributed computing
- **Contact stanford@nist.gov if you are interested in a working group**

Slide 26

GRID COMPUTING INFRASTRUCTURE

Geoff Brown
Oracle Corporation
Redwood Shores, CA

THE IT CHALLENGE

Your IT department is under constant pressure. You have to implement, maintain and improve the operational systems that run your companies and also to design and create additional systems that can provide competitive advantages for your business. These systems could be used for a variety of purposes, from deeper analysis of market and business trends, to an improved customer service experience, to reducing overall product costs.

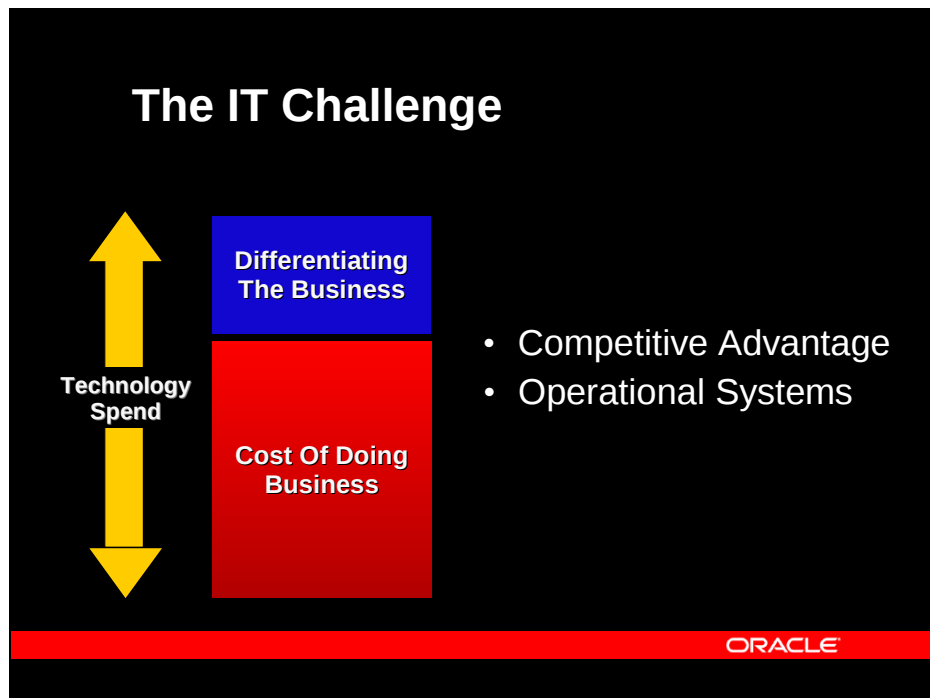


Figure 1

THE IT RESPONSE

To succeed, you must try to meet this challenge, but also deliver even greater value. One effective way to accomplish this is to reduce the portion of the budget required to meet operational costs, which allows you to use more of your resources to provide competitive advantage. And you would like to accomplish this while saving both money and time to market.

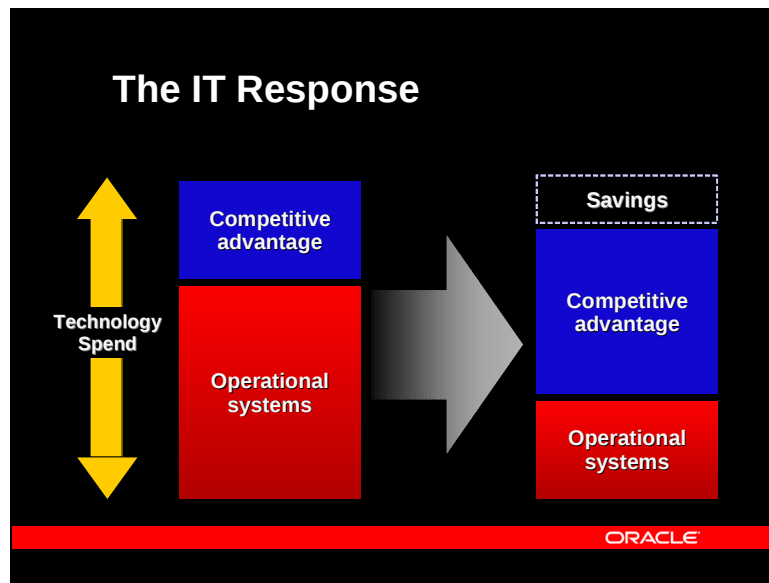


Figure 2

Infrastructure Costs

- Low Utilization of Processor Resources
- Low Utilization of Storage Resources
- Weak Systems Management Capabilities
- Weak Asset Management Capabilities
- High Cost
- Slow Provisioning
- Inadequate SLA's

...and It is Only Getting Worse!

ORACLE

Figure 3

ROBUST AND FLEXIBLE INFRASTRUCTURE

You can reduce your overall costs with a powerful, robust and flexible infrastructure. Rather than having the complexity of your infrastructure be a budget consumer, taking time and resources away from your overall budget, you could choose the right infrastructure and have your choice reduce your total cost of ownership while increasing the productivity of all of your IT staff. As Nick Gall of META Group says, the better the infrastructure, the greater the benefits.

The main way an infrastructure can provide value is through increased productivity. The more functionality your infrastructure supplies, the less time you will have to spend implementing and maintaining that functionality in your IT systems. Providing productivity is half the equation. If you have to spend the same amount of time implementing a feature in your infrastructure that you would in your application systems, the net benefit is zero. The easier it is to obtain a benefit, the greater the overall value.

Another important aspect of a standardized IT infrastructure is that you can use it over and over without any additional implementation work. An infrastructure that can provide ongoing benefits from your original investment will provide the greatest value for your organization.



Figure 4

STANDARDIZED INFRASTRUCTURE ELEMENTS

To achieve these goals, your business needs an Unbreakable Software Infrastructure. Oracle's Unbreakable Software Infrastructure provides a wealth of functionality. The advanced features of Oracle's Unbreakable Software Infrastructure can help to solve your tactical business problem today, as well act as a strategic investment in the future of all your IT systems.

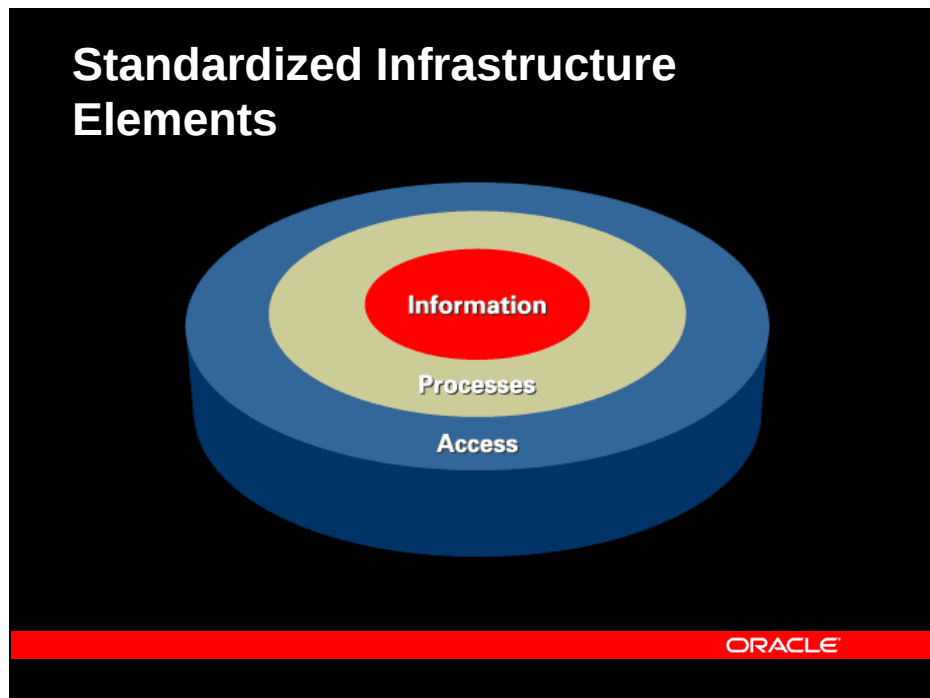


Figure 5

STANDARD SOFTWARE INFRASTRUCTURE

The core of the Oracle Unbreakable Software Infrastructure is information - your data. Your data is one of the most valuable resources of your company. Virtually all of your information systems are built on your data. By making data the core of an Unbreakable Software Infrastructure, you are building on the core of your company's valuable information.

With Oracle's Unbreakable Software Infrastructure, you can keep all of your data in one centralized repository - data from your transactional (OLTP) systems, data used for business intelligence functions, and a wide variety of other documents, such as Web content, E-mail, and calendar and resource scheduling information.

This centralized repository reduces your overall management overhead and allows you to consolidate the number of servers in your organization, which will further reduce overhead. By having all of your data in a single repository, you also reduce the need for resource-consuming data transfers required for multiple uses of the same data.

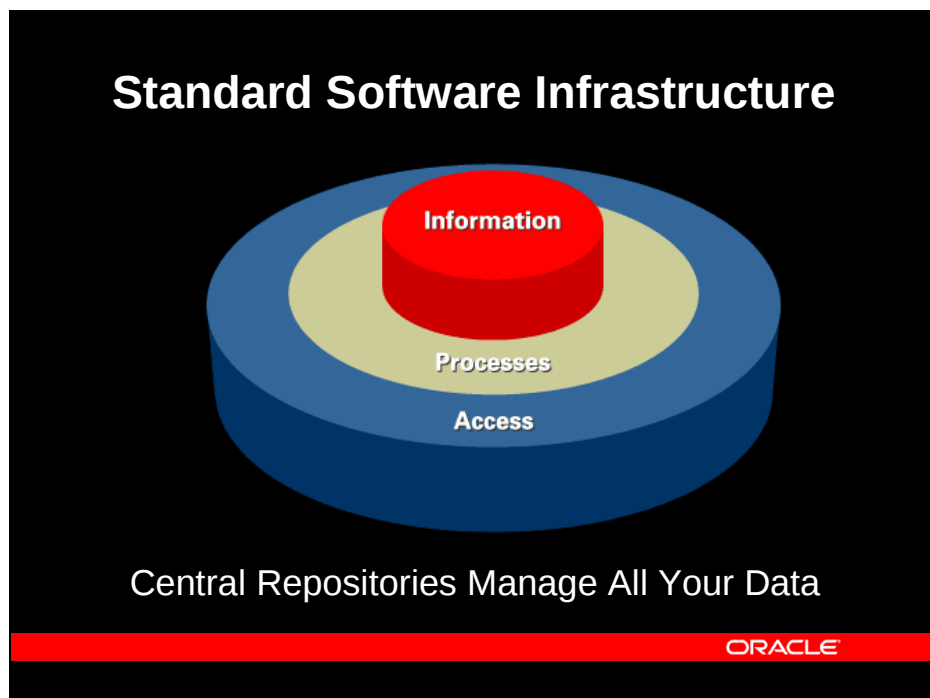


Figure 6

BENEFITS OF CENTRALIZING DATA

The advantages built into Oracle's Unbreakable Software Infrastructure provide transparent benefits for all your data. For instance, Oracle's Infrastructure gives access to all data for large numbers of users, without any performance impediments caused by locking issues or extensive coding to work around potential problems. Oracle's Infrastructure can scale up or out, seamlessly, which guarantees you both scalability and the cost benefits of using commodity hardware. Oracle lets you add the widest variety of indexes to all of your data, which can in turn provide rapid access. You can even define your own custom indexes for your own specific data. Oracle9i gives you a new feature which will automatically compress the stored representation of your data, saving you storage space and improving the performance of your application systems and maintenance operations. And Oracle's Unbreakable Software Infrastructure lets you separate any and all of your data into partitions in many different ways – for maintenance, security or performance considerations. These benefits are all available to all your applications – without any additional coding or maintenance on your part. The cost of these benefits is zero.

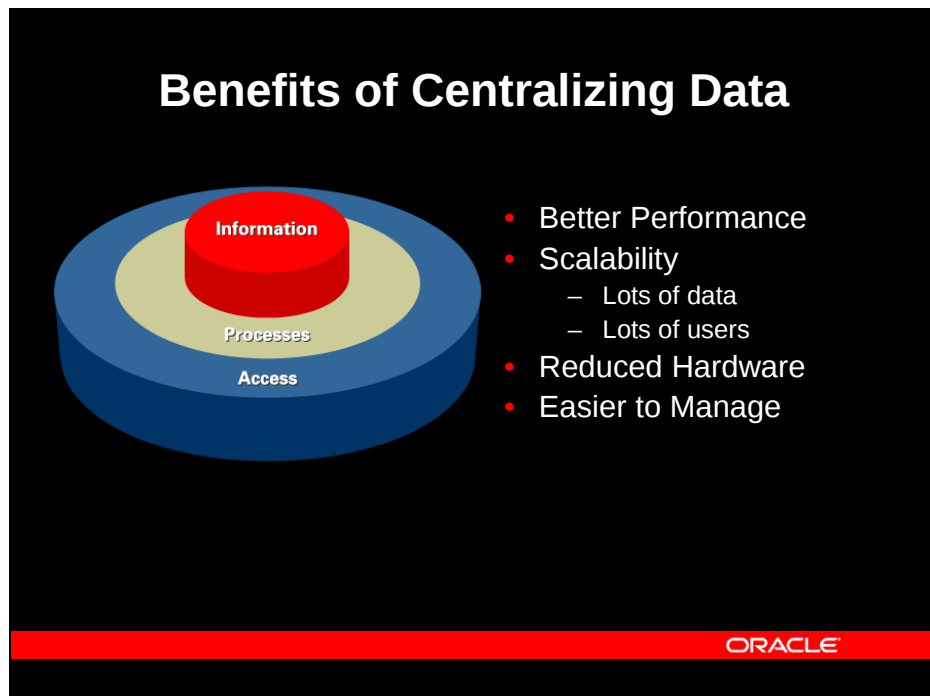


Figure 7

STANDARD SOFTWARE INFRASTRUCTURE

Data is the core of our Unbreakable Software Infrastructure, but information systems do more than simply store and retrieve data. Systems are built to interact with data to create business processes used to support and enhance business operations.

Oracle's Unbreakable Software Infrastructure helps to create and deploy your business processes. The advantages of Oracle's Infrastructure help you to create business processes quickly and efficiently to respond to the demands of your environment.

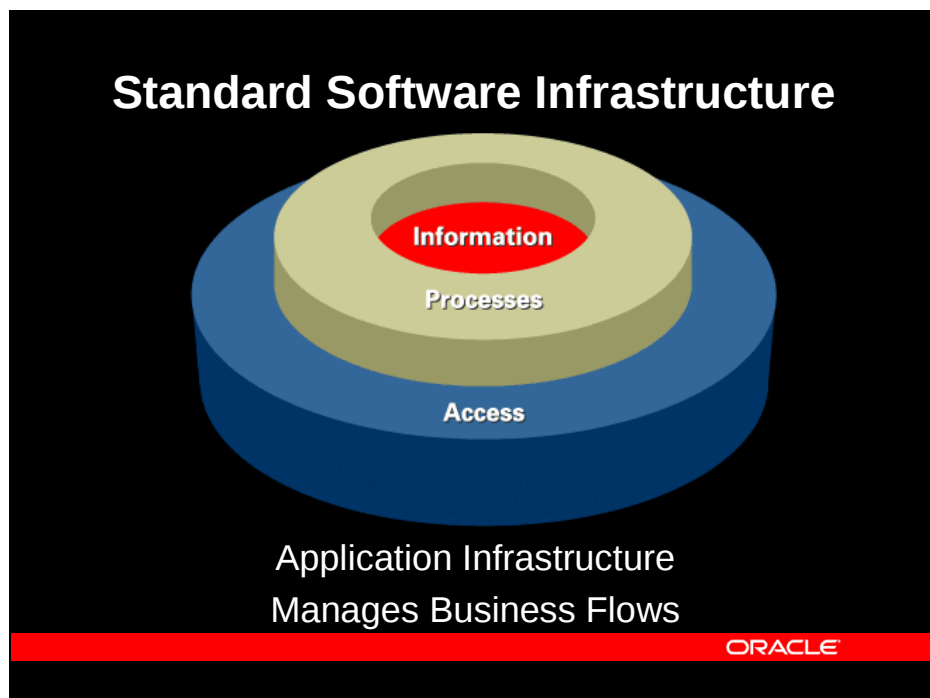


Figure 8

BENEFITS OF STANDARDIZED APPLICATION ARCHITECTURE

Oracle's Unbreakable Software Infrastructure includes features that can make it easier for you to create and maintain your vital business processes. You can create data-aware Java components, pre-baked with all the functionality they will need to access and manipulate data.

Our Infrastructure includes transparent caching for not only data, but the results of processes, such as HTML pages or fragments. Retrieving cached data is much faster than recreating it, and your application systems will perform better – transparently. Oracle includes tools to easily manage the way you use this caching.

Oracle includes a special feature to pre-calculate aggregate values, which are frequently used in data warehousing. Of course, you can use this capability without any modifications of any of your applications. Oracle's Infrastructure even includes wizards to suggest which pieces of data could benefit from this type of pre-calculation.

Your own business processes are unique to your own business situation. That's why Oracle lets you create your own functions, which you can use in any application or SQL code, just like standard built-in functions. The productivity gains provided by the Unbreakable Software Infrastructure can extend into the particulars of your own specific business.

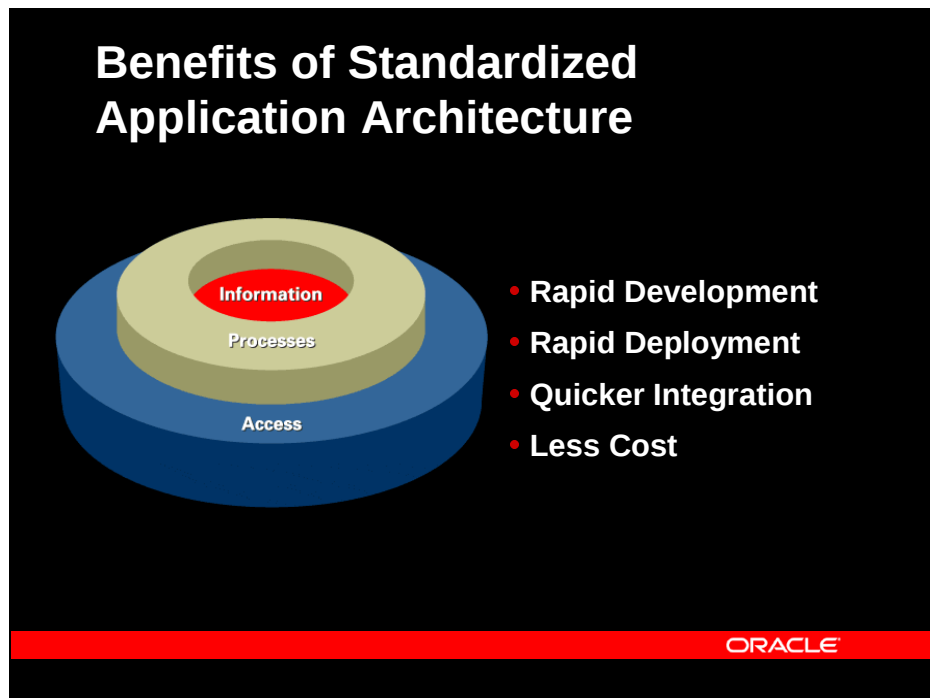


Figure 9

STANDARD SOFTWARE INFRASTRUCTURE

Processing is not the end goal of any information system. You also have to deliver the results of those business processes to your clients across a wide range of channels. Using a single business process across many channels can significantly improve the productivity of your development effort, as well as reduce the need for redundant systems that require constant maintenance and synchronization.

The final step of data access is also where an infrastructure becomes truly Unbreakable with efficient and powerful security mechanisms.

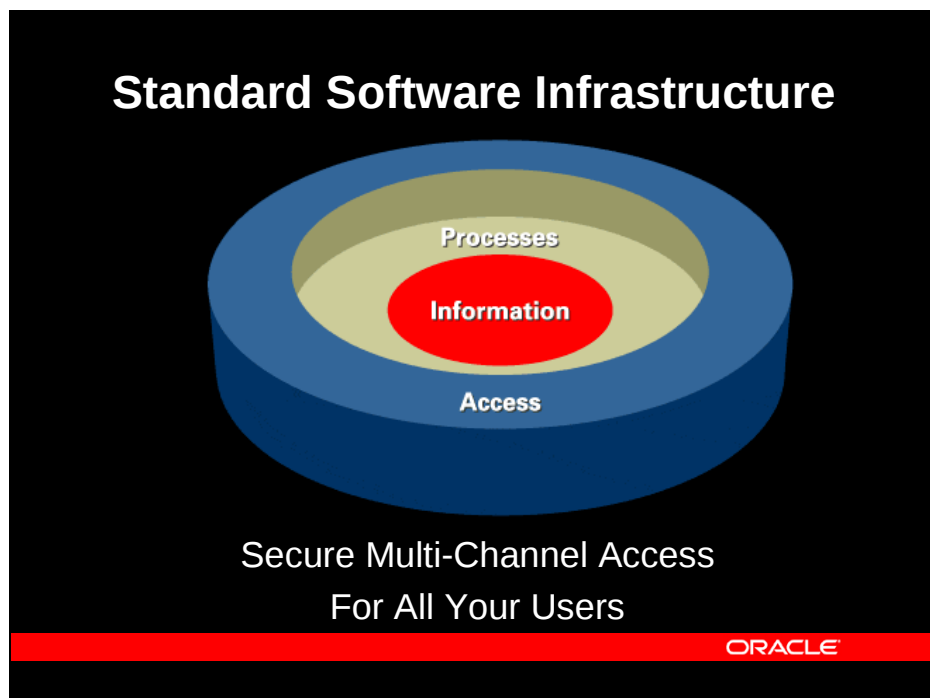


Figure 10

BENEFITS OF MULTI-CHANNEL ACCESS

You can propagate the results of your business processes to multiple channels, without having additional logic or redundant applications to address the needs of each individual channel. Whether the final destination of the information generated by a process is a standard client machine, a Web page, a portal or a mobile device, Oracle's Infrastructure provides easy support for each channel.

To make it easier for your users to access the data they need, Oracle's Infrastructure provides powerful search capabilities.

Oracle has been a leader in secure access for many years. Built into our Unbreakable Software Infrastructure are features that can provide a single digital identity for all applications, so that your users only have to log on once a day.

Oracle has extremely flexible security, which allows you to limit access to data based on the value of the data. For instance, one column in a table could have a value that is used as a label to allow or prevent access to the information in that row. You can implement this content-based security on the data itself, so it will apply for all systems that access the data.

Oracle provides encryption of your data as it is stored and in transit, as well as selective encryption if needed.

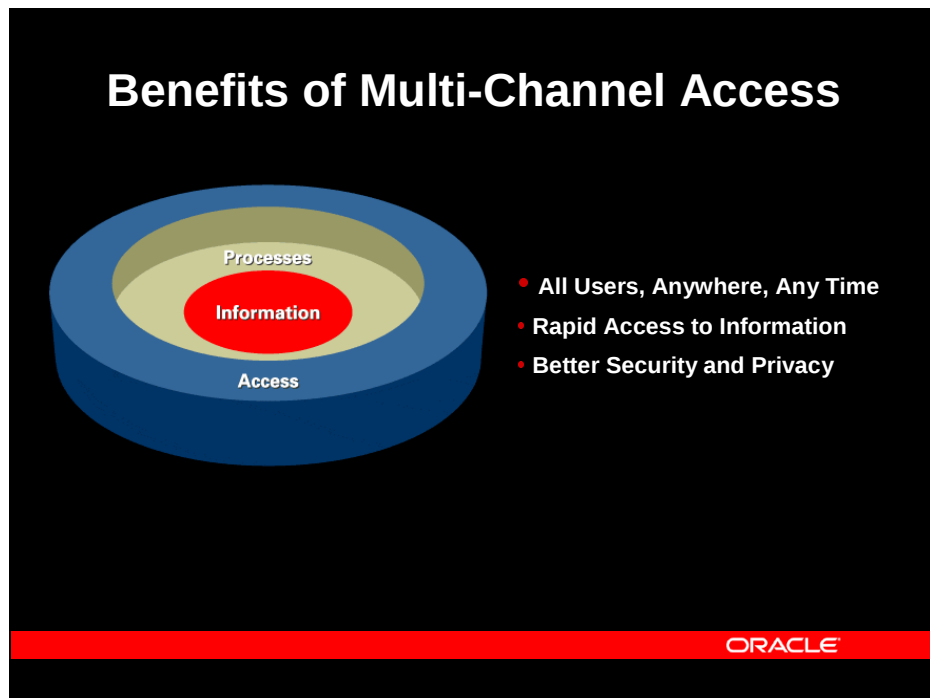


Figure 11

UNBREAKABLE SOFTWARE INFRASTRUCTURE

Oracle's Unbreakable Software Infrastructure excels in 5 areas crucial to the value of any infrastructure – performance, scalability, availability, security and manageability. You will be seeing examples of customers and independent proof points for Oracle's leadership in each of these areas throughout the day today.

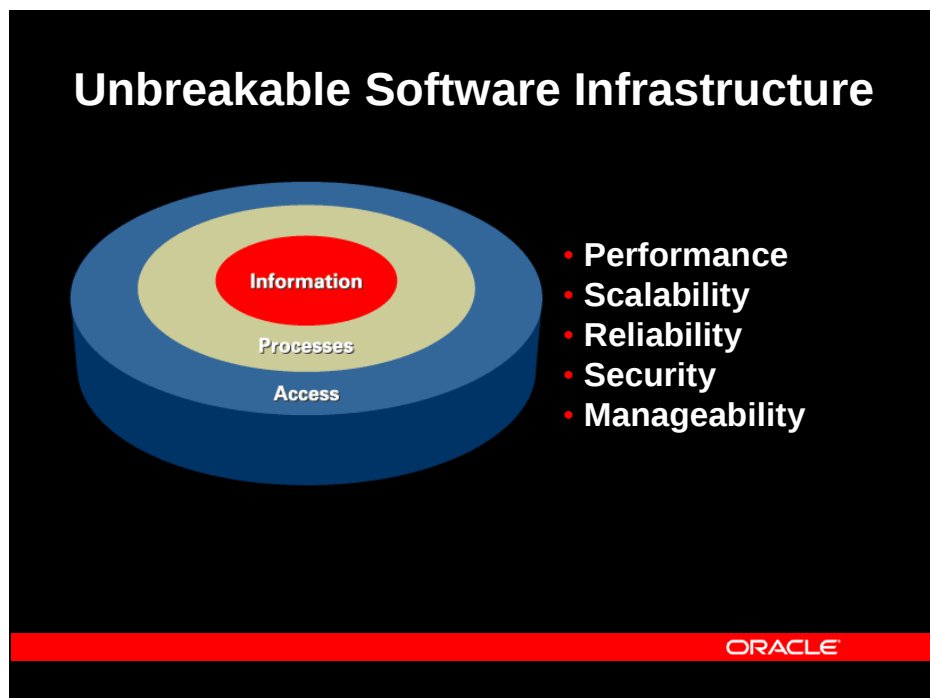


Figure 12

EVOLUTION OF COMPUTING

Of course there have been other IT infrastructures in the past. In the Age of Big Iron customers used mainframes to run their infrastructure. These had significant advantages in quality of service and efficiency. But they were also inflexible, leading to large application backlogs, and costly.

Client-Server computing arose in response to this. This swung the pendulum to the opposite pole by highly distributing systems. While this reduced the initial purchase price of systems and provided greater flexibility, it also cost more in integration and quality of service problems.

Next generation infrastructures balance these centralized-decentralized designs by gaining the advantages of consolidation while retaining flexibility in application design.

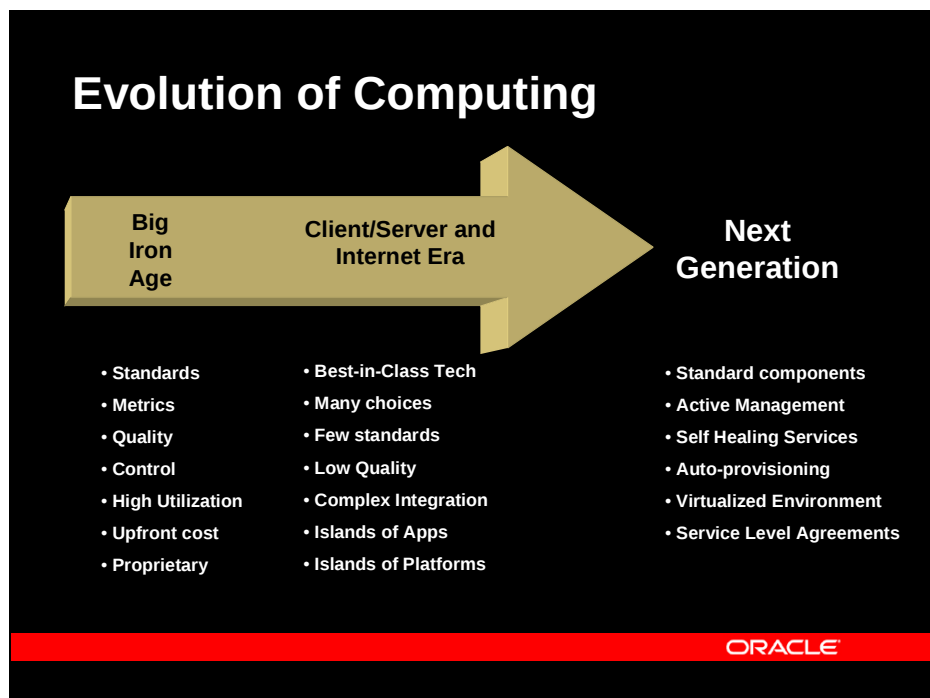


Figure 13

NEW INFRASTRUCTURE DESIGN STRATEGY

This new infrastructure model is based on a new design strategy based on a few simple principles. Build a single infrastructure for your entire IT department. Think holistically about data management, application processing, networking and user access. Create a few large pools of resources that can be used across all applications, not individual islands of systems.

Standardize on a few pieces of infrastructure software – databases, application servers, etc. Best of breed technology is not cost-effective. Enforce the adoption of architecture standards for all your applications. Make your applications take full advantage of the standardized infrastructure over time. Ensure that you have a comprehensive end-to-end system management solution for your infrastructure. You cannot scale up your infrastructure without resolving this issue.

Grid Computing

Persistent environments that enable software applications to integrate instruments, displays, computational and information resources that are managed by diverse organizations in widespread locations.

-- The Globus Project

ORACLE

Figure 14

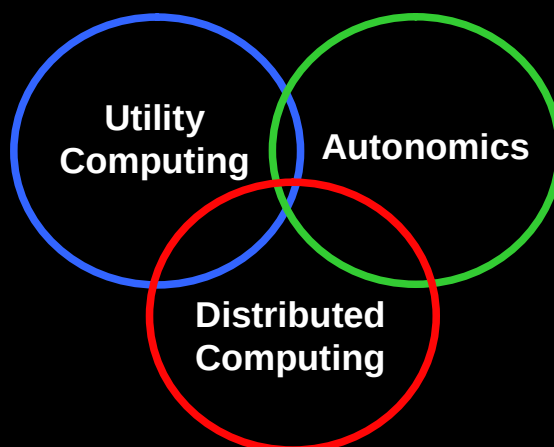
New Infrastructure Design Strategy

- Attack the Whole Problem
 - Data, Applications, Users
- Virtualize Resources for Flexibility
- Standardize Infrastructure Software
 - Best of Breed is not necessarily most cost-effective
- Enforce Application Architecture Standards
- Comprehensive System Management

ORACLE

Figure 15

Infrastructure Design Elements

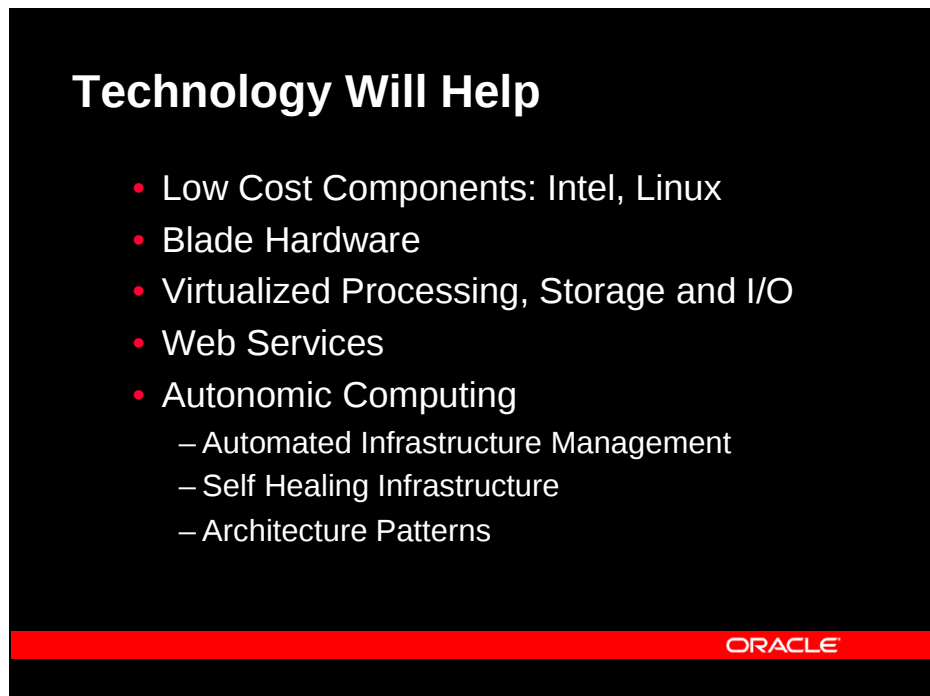


ORACLE

Figure 16

TECHNOLOGY WILL HELP

There are numerous technical breakthroughs which have made this possible including low cost computing components, new clustering designs that enable a modular approach to system design; web services and other integration techniques; and improvements in self-managing computing systems that enable greater scale up.



Technology Will Help

- Low Cost Components: Intel, Linux
- Blade Hardware
- Virtualized Processing, Storage and I/O
- Web Services
- Autonomic Computing
 - Automated Infrastructure Management
 - Self Healing Infrastructure
 - Architecture Patterns

ORACLE

Figure 17

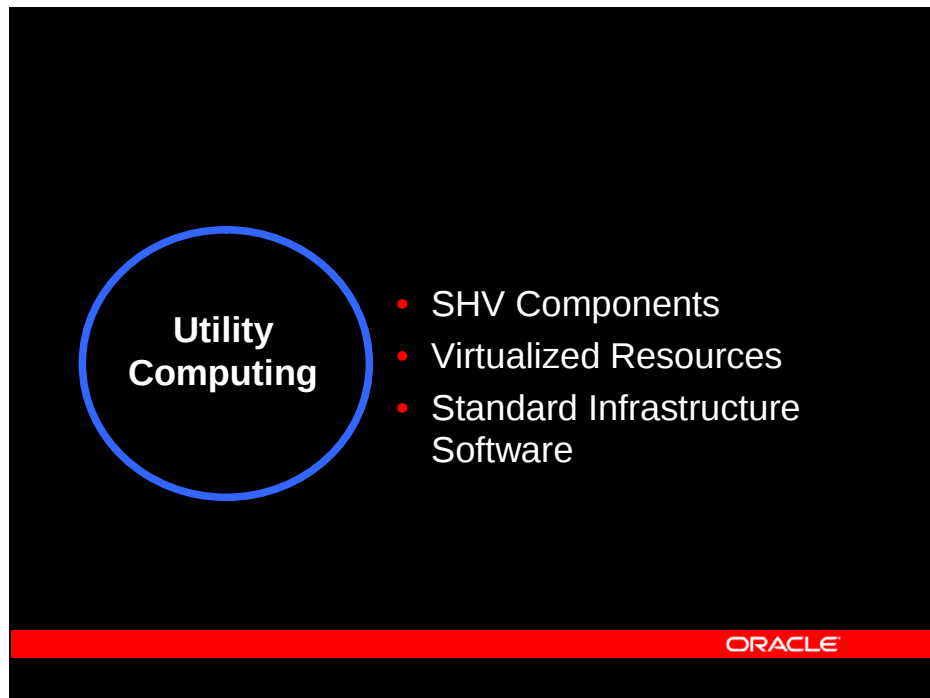


Figure 18

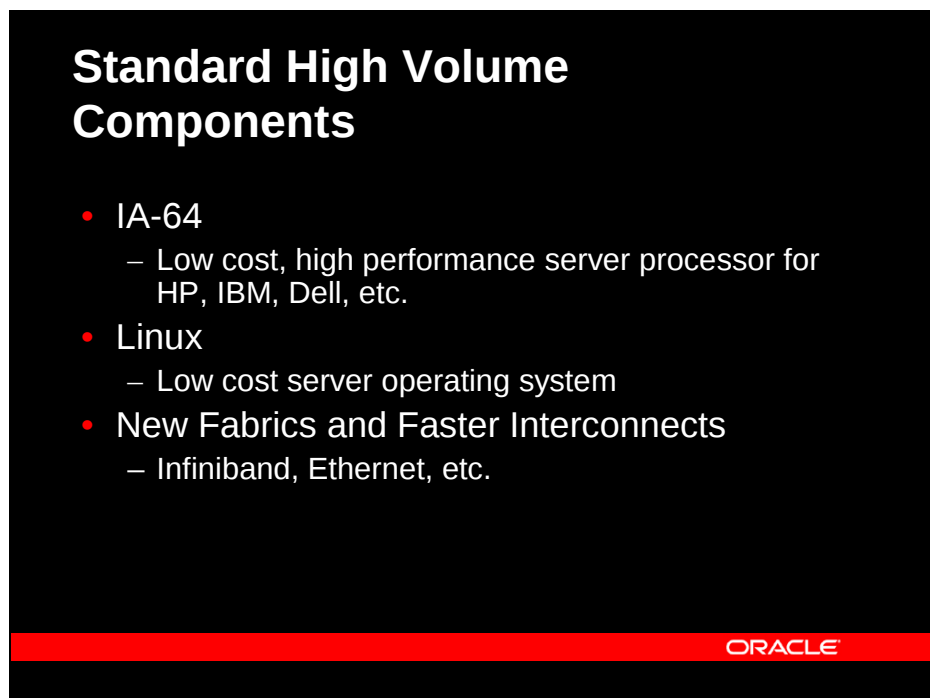


Figure 19

Virtualized Data Center Resources

- Centralized Pool of Resources
 - Storage
 - Processing
 - I/O
- Resources can be Isolated and Dedicated

ORACLE

Figure 20

What's the Problem with Storage?

- Islands of Storage
- Storage Tightly Coupled to Applications and Servers
- Storage Utilization is often < 50%
- Storage Administration Costs are Sky Rocketing

...Storage is growing 30%+ per year

ORACLE

Figure 21

STORAGE VIRTUALIZATION

With storage virtualization we eliminate the islands of storage. By consolidating and virtualizing we can dramatically reduce waste and inefficiency.

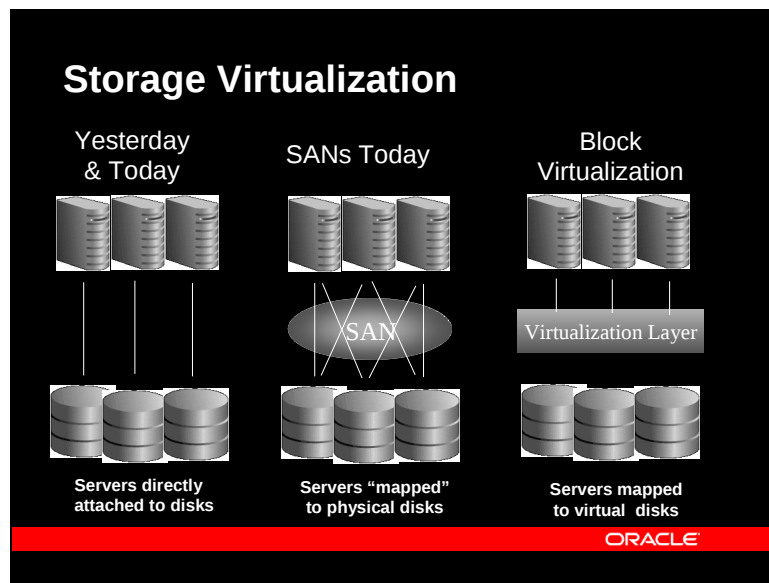


Figure 22

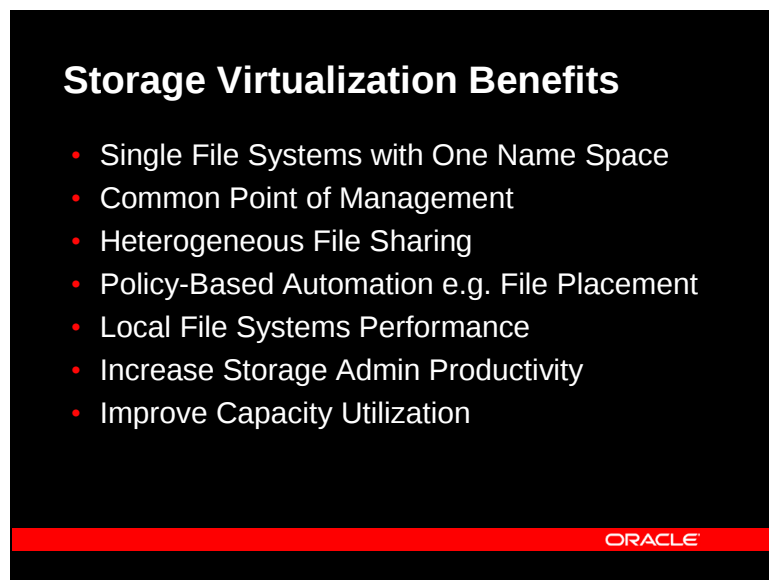


Figure 23

PROCESSOR VIRTUALIZATION

Through these blades we can create a huge pool of computing capacity available on demand. There are also new partitioning capabilities being built into SMP systems that enable sharing of resources. This is another approach that appeals to customers.

What's the Problem with Servers?

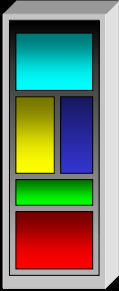
- Islands of Servers for Each Application
- Too Many Independent Servers
- Average CPU Utilization is Low (< 25%)
- Multiple Unique Vendors and Versions
- Complex and Different Software Stacks
- Poor Systems Management
- Slow Provisioning

... the More Servers Added, the Lower the Utilization of Assets.

ORACLE

Figure 24

Processor Virtualization

Blade Servers	Virtual Partitioning Servers
	
- Hundreds of Intel processors in a single rack. - Sophisticated management tools - Self-healing capabilities - Excellent for Web Servers, Clusters, etc.	- Large multi-processor systems. - Physical and virtual partitioning - Excellent for consolidating servers - Dynamic CPU utilization

ORACLE

Figure 25

VIRTUALIZING THE DATA CENTER

This diagram shows how a data center moves from separate islands of resources today to virtualized pools of resources.

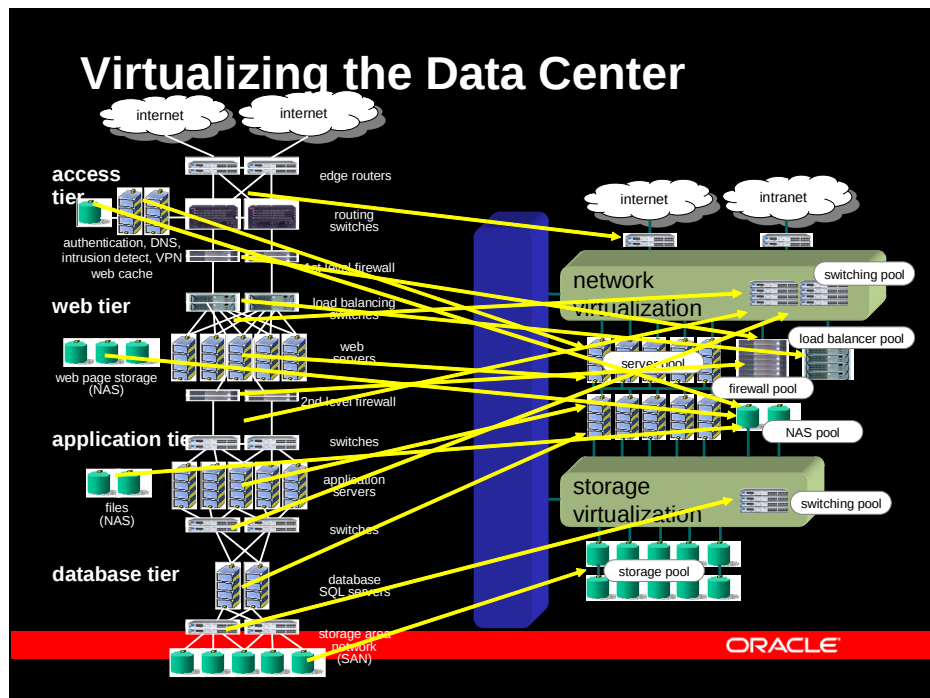


Figure 26

UTILITY COMPUTING EXAMPLES

CSFB: The BladeFrame gives CSFB flexibility. Provisioning and configuration tasks that require three weeks or more with legacy servers are performed in just minutes on the Egenera system, enabling us to accommodate growth and launch new applications in a timeframe never before possible. Simplifying server deployment also allows developers to focus on strategic initiatives, which means they can respond more quickly to business opportunities. The BladeFrame helps CSFB adapt to change faster, giving them a powerful advantage in a highly competitive market.

CDC IXIS Capital Markets: The financial modeling CDC IXIS Capital Markets uses to predict the outcomes of particular events provides the entire underpinning of their operation. They use a sophisticated application that is able to model a number of parameters and help them to predict the most likely outcome of a position. The key is to ensure that every aspect of what might happen in each market is covered, and to be able to report the effects of changing parameters rapidly to the customer. To handle that volume of complex modeling calls for a great deal of processing power, but no storage. CDC IXIS Capital Markets uses a server farm with NEBS Level-3 certified Netra[tm] t1 Model 100/105 servers from Sun Microsystems, Inc. to run the financial modeling application.

Utility Computing Examples

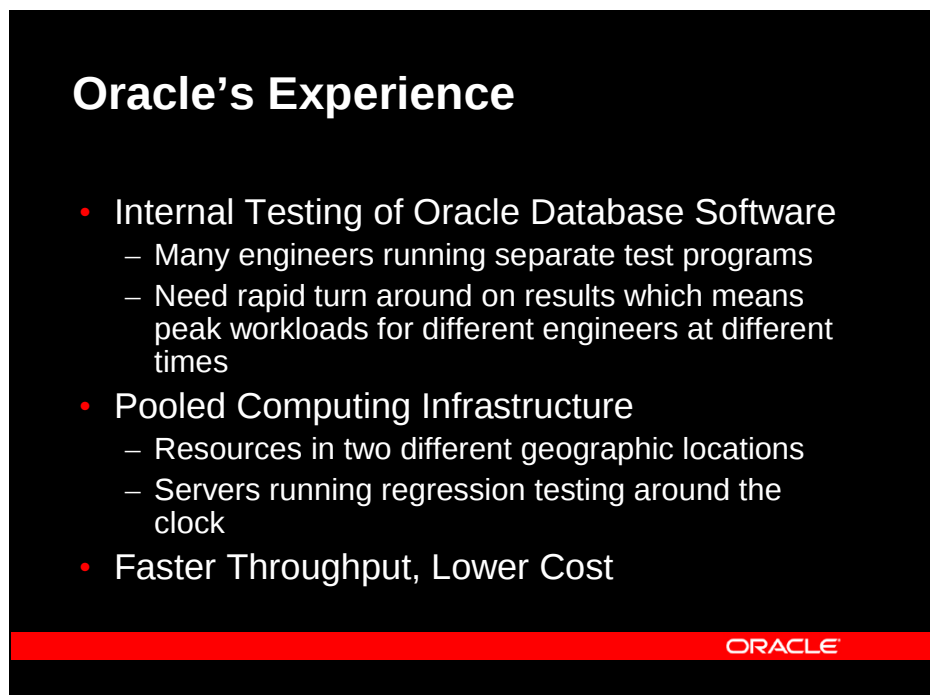
- Oracle
 - Solaris blades automating RDBMS development
- Credit Suisse First Boston
 - Egenera blade farm supporting financial and web applications
- Department of Energy
 - 1,400 Linux blades to study materials design
- Celera
 - 1000 Alpha nodes analyzing human genomic data

ORACLE

Figure 27

ORACLE'S EXPERIENCE

Oracle has also moved to a Utility Computing model for some of its internal application development

A presentation slide with a black background and white text. The title 'Oracle's Experience' is at the top left. Below it is a bulleted list with three main items: 'Internal Testing of Oracle Database Software', 'Pooled Computing Infrastructure', and 'Faster Throughput, Lower Cost'. Each item has sub-points. At the bottom right, there is a red horizontal bar with the word 'ORACLE' in white capital letters.

Oracle's Experience

- Internal Testing of Oracle Database Software
 - Many engineers running separate test programs
 - Need rapid turn around on results which means peak workloads for different engineers at different times
- Pooled Computing Infrastructure
 - Resources in two different geographic locations
 - Servers running regression testing around the clock
- Faster Throughput, Lower Cost

ORACLE

Figure 28

AUTONOMICS

Even as companies move to a utility computing model, the issue of resource management becomes critical. As you build larger and larger computing pools, human beings become challenged to manually manage these resources effectively. The answer is to build systems that manage themselves.

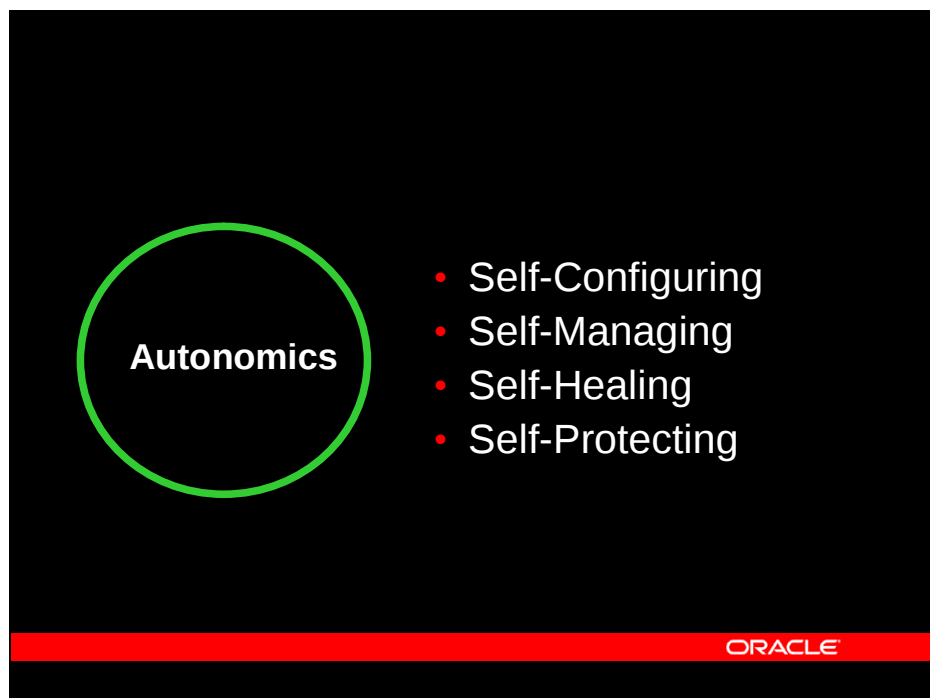


Figure 29

WHAT'S THE PROBLEM WITH SOFTWARE INFRASTRUCTURE?

So we've talked about the movement to low cost, high volume hardware and then pooling this hardware to create a large flexible resource for all your applications. The next problem is the software infrastructure. Today most customers use multi-vendor solutions. This complexity diverts focus from the business requirements

Data is ignored

Politics rule

The IS organization and business users can't work together

There is no plan

Processes are implemented for the enterprise, not the customer

A flawed process is automated

No attention is paid to skill sets

The key is to standardize the infrastructure software so that there is an easy way to install, maintain and upgrade your infrastructure. Applications can be written to use this infrastructure in a consistent fashion. All this means gains in efficiency.

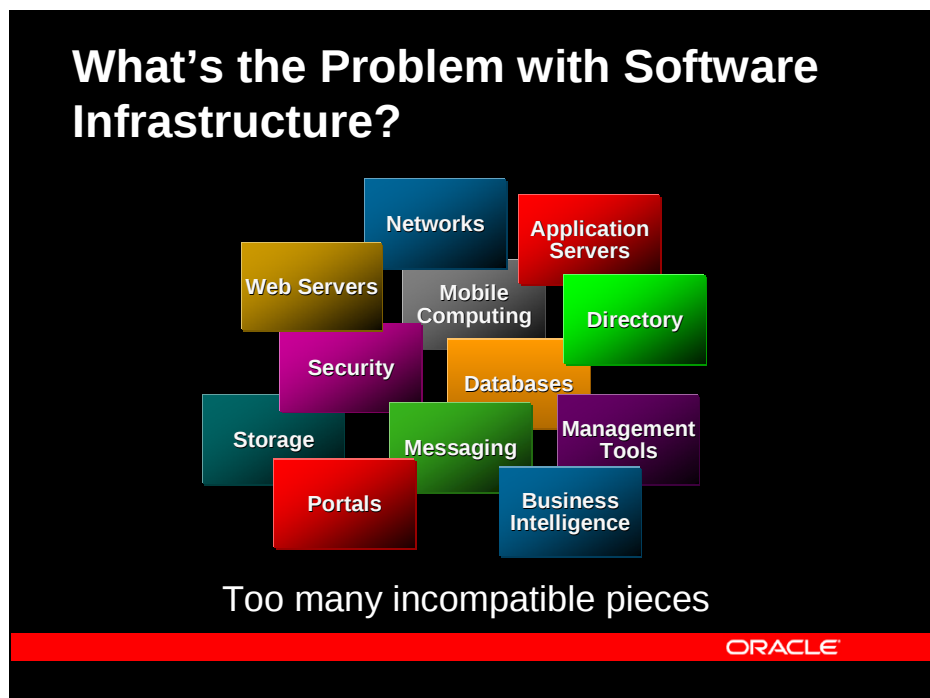


Figure 30

WHAT'S THE PROBLEM WITH SYSTEM MANAGEMENT?

Expensive to operate

Manual labor intensive deployment and changes
Inefficient asset utilization
Weak systems mgmt

Inflexible and complex

Multiple architectures for apps and customers
Highly complex because they are all different
Integration is complex and costly

Error prone, unreliable and slow

Human factor in change requests
Limited high availability built for specific apps only
Lack integrated management

...how are we going to manage our environment in the next decade?

What's the Problem with System Management?

- Expensive to Operate
- Inflexible and Complex
- Error Prone, Unreliable and Slow

... How are We Going to Manage Better?

ORACLE

Figure 31

CHALLENGE: WORKLOAD MANAGEMENT

The second challenge is how to deal with the allocation of resources within the pool to various applications. It's not just meeting the processing demands of the applications but also determining appropriate HA strategies for each application, dealing with spikes in demand that randomly occur, as well as figuring out future capacity needs.

This can be far beyond the ability of humans to handle when you're dealing with a huge resource pool.

Challenge: Workload Management

- Installation, Configuration, Backup/Restore
- Partition/Control Short-Term Load Among Nodes
- Allocation of Nodes
 - HA spares
 - Handling spikes
 - Integrate capacity planning w/ growth as needed
- Support for Necessary User Choice
 - Multiple OS in the same rack (Windows, Linux, Proprietary UNIX)

ORACLE

Figure 32

SOLUTION: END-TO-END SERVICE LEVEL MANAGEMENT

So the goal of autonomies is to enable the infrastructure to manage itself with minimal intervention by human beings. People will still set high-level business and technical policies about the infrastructure but the system itself will do installations, maintenance, tuning, recommend capacity plans and so on. All of this will be reported through comprehensive graphical displays.

Single system image for competing workloads running within multiple server farms
Automated management of workload groups in response to service metrics
One event system across and between stateless and stateful cluster domains. One big happy cluster: no mid-tier/backend distinction. Mixed storage, flexible mapping services to nodes

I'm not going to talk about customer successes in this section. All customers are taking advantage of autonomies to some degree today. Autonomics have been incorporated into software for some years and will continue to be refined in coming years. Everyone is and will continue to use these capabilities without really having to know much about it.

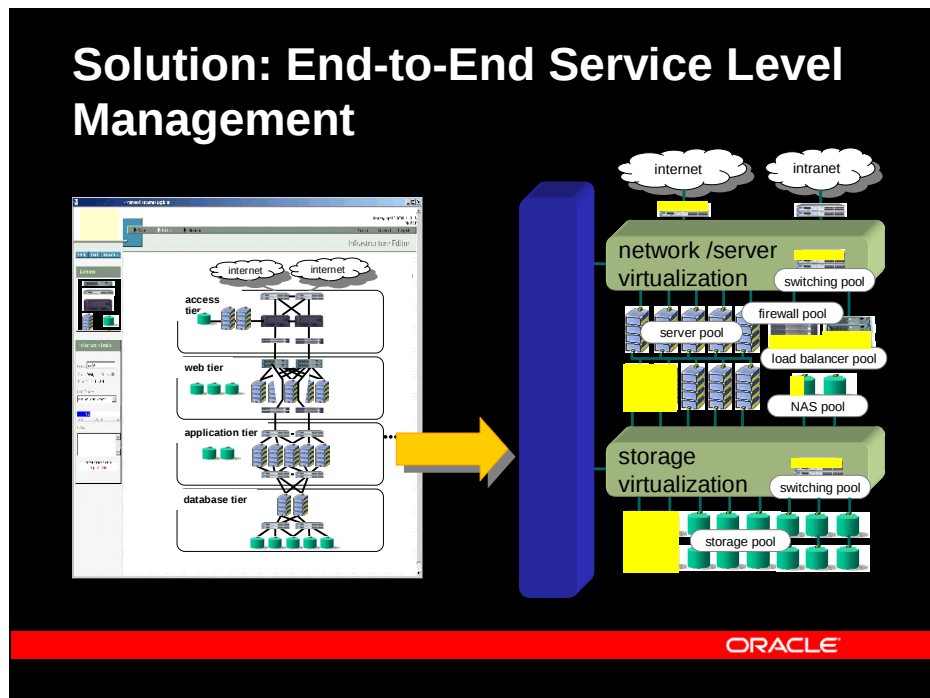


Figure 33

DISTRIBUTED COMPUTING

The final aspect of this new infrastructure is Grid Computing. Grid computing is a set of technologies that take into account the fact that not all applications and all data will necessarily reside in a single resource pool. There are many times when resources are distributed and applications and users must access these resources in a distributed fashion.

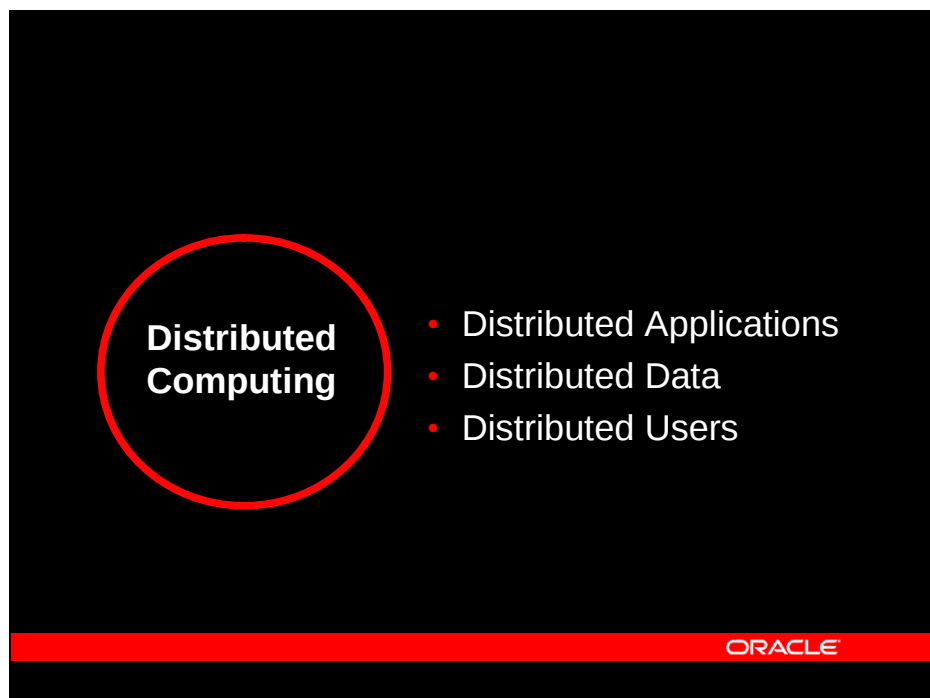


Figure 34

DEFINING DISTRIBUTED COMPUTING

Grid computing provides standard interfaces between resource pools so that users and applications in one location can access resources in another location transparently.

The Grid is not an alternative to the Internet. It is a set of additional protocols and services that build on Internet protocols and services to support the creation and use of computation and data-enriched environments.

Grid computing is not web services. Web services are one technology for implementing distributed applications, but there are many more technologies involved in Grid computing.

Grid computing is not peer-to-peer computing. P2P is one style of distributed or Grid computing that relies on systems in place for a limited class of parallel processing applications.

Grid computing is also not the same as application hosting. Application hosts may use Grid Computing technology to make their resources available, but app hosting is more about the business issue of outsourcing work than about distributed computing technology.

Defining Distributed Computing

- Network of Clients and Service Providers
 - Standard interfaces and universal availability
 - Resource sharing, fault tolerance, and load balancing
- What It's Not
 - Next Generation Internet
 - Just Web Services
 - Just Peer-to-Peer Computing
 - Application Hosting

ORACLE

Figure 35

GRID ARCHITECTURE

A Grid architecture consists of different applications using a meta-reservation service to schedule run-time. The scheduler uses a resource management protocol to identify appropriate computing resources around the network to run the applications. The computing resources need to advertise their characteristics such as types of applications they can run, databases they have access to, etc.

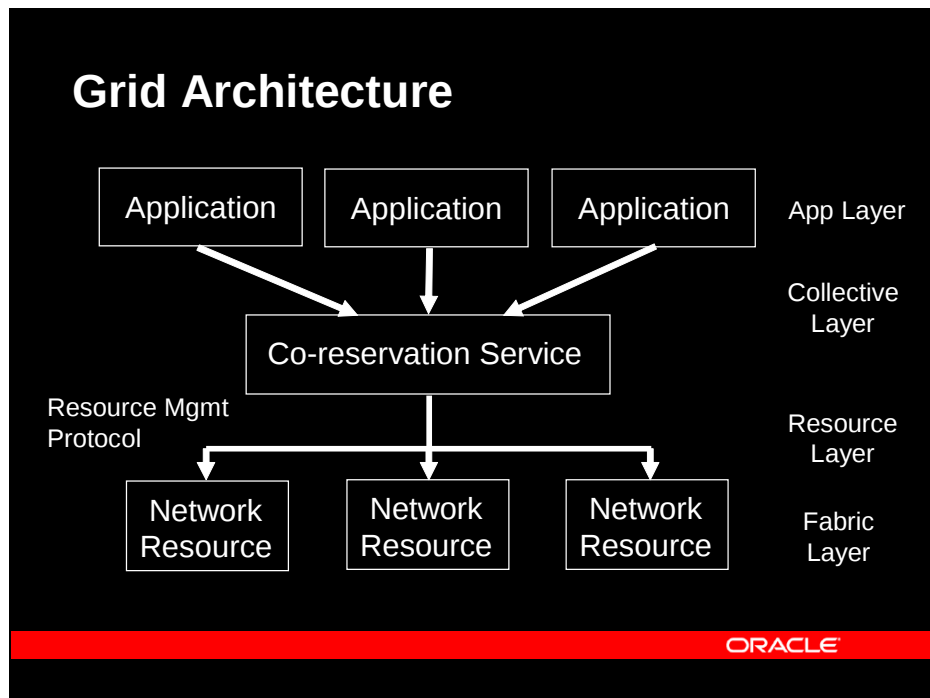


Figure 36

WHY SHOULD YOU CARE?

There are some important benefits of Grid computing. You can now create gigantic applications that otherwise would not have the resources to run. You can access data that would otherwise be unavailable.

This is extremely valuable for many day-to-day applications. Many companies need regular access to remote data sources such as consumer credit information, medical or genomic databases, and more. These can now be made a transparent part of any application that is granted access. This means more real-time availability of information and better applications.

Grid computing can also solve problems with mergers and acquisitions. When a company spins off a subsidiary it is often difficult to immediately separate ERP or CRM systems. Grid computing can provide access to existing company applications to the former subsidiary while maintaining real separation of information, proper billing for usage, etc.

Grid computing can also be attractive for companies that have many suppliers or distributors that need to share information about manufacturing processes, inventories, etc. By having appropriate access to the data warehousing information from your supply or distribution chain you can do much better data mining or simulation work.

Why Should You Care

- Completely Changes the Economics of Computing
 - Drastically lowers cost
 - Extends availability of resources
- Valuable for Many Organizations
 - Accessing external data sources – credit bureaus, genomic or medical databases, etc.
 - Mergers & acquisition situations
 - Large supply or distribution chains

ORACLE

Figure 37

STANDARDS EFFORTS

Grid computing is gaining momentum.

Standards Efforts

- Global Grid Forum
 - Provide standards specifications for grid technologies
 - Comprises over 200 organizations
 - Oracle UK co-chairs the *Data Access and Integration* working group
- Globus
 - provides open source toolkit conforming to Global Grid Forum specifications

ORACLE

Figure 38

DISTRIBUTED COMPUTING CUSTOMER EXAMPLES

Gene Logic chose the AVAKI technology in order to maximize the use of its existing internal computing infrastructure, which has cyclical usage characteristics. The AVAKI technology will be employed to create more efficient utilization of this infrastructure, and more rapid completion of certain internal analysis efforts.

Deutsche Bank is one of the leading international financial service providers. Its investment banking division relies heavily on technology to meet the computing needs of its traders around the world. While traders at Deutsche Bank's New York office use high-end Pentium desktop and UNIX workstations to effectively carry out their daily functions, end of day reporting and analysis, the need for optimal compute power intensifies. Following in the footsteps of its colleagues in Frankfurt, Germany, the NY office implemented Platform's workload management solution, Platform LSF. This allowed Deutsche Bank to create a virtual mainframe from its existing cluster of computers, and eliminated the need to purchase additional hardware to address their demanding computing needs.

European Aeronautic, Defense and Space Company (EADS) is Europe's largest aerospace company, resulting from the combination of Aerospatiale Matra SA and DaimlerChrysler Aerospace AG (DASA). Prior to the merger between the two companies, DASA's computing demands had traditionally been met by central mainframes, with some peak requirements satisfied by external supercomputers. Over a period of time, this environment had been replaced by distributed workstation and server systems. In an effort to replicate the easy-to-use, centralized mainframe environment, DASA adopted Platform's workload management solution, Platform LSF MultiCluster, to manage their computing workload and distribute batch jobs across the network to the most suitable computers.

GriPhyN: Communities of thousands of scientists, distributed globally and served by networks of varying bandwidths, need to extract small signals from enormous backgrounds via computationally demanding analyses of datasets that will grow from the 100 Terabyte to the 100 Petabyte scale over the next decade. The computing and storage resources required will be distributed, for both technical and strategic reasons, across national centers, regional centers, university computing centers, and individual desktops

Distributed Computing Customer Examples

- Scientific: European Data Grid (CERN)
 - Thousands of physicists analyzing petabytes of distributed elementary particle data
- Aerospace: DaimlerChrysler Aerospace
 - Use distributed servers to perform complex simulations

ORACLE

Figure 39

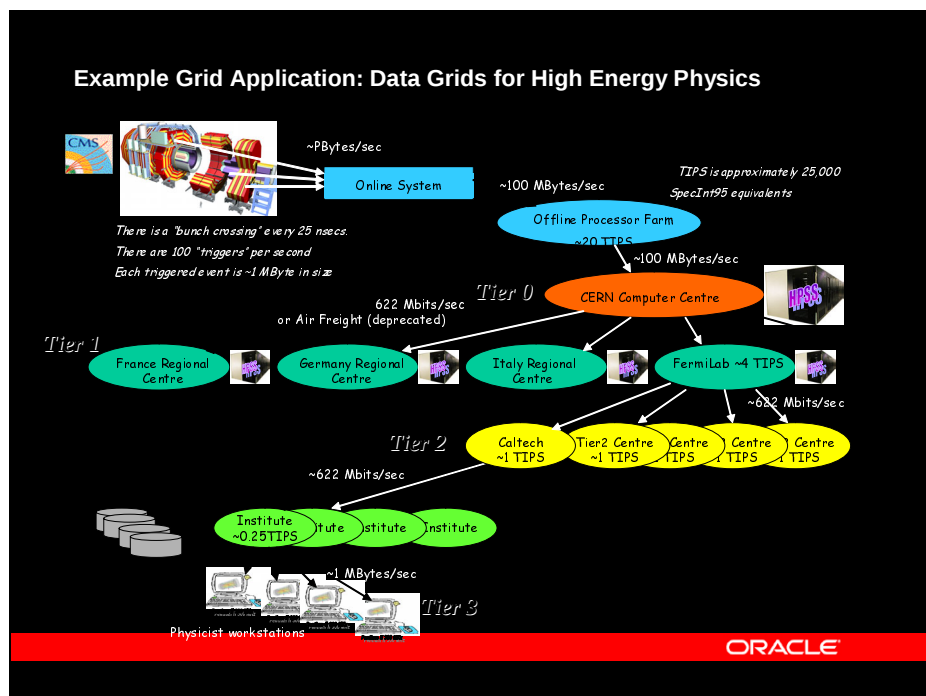


Figure 40

Example Grid Application: Stanford Linear Accelerator

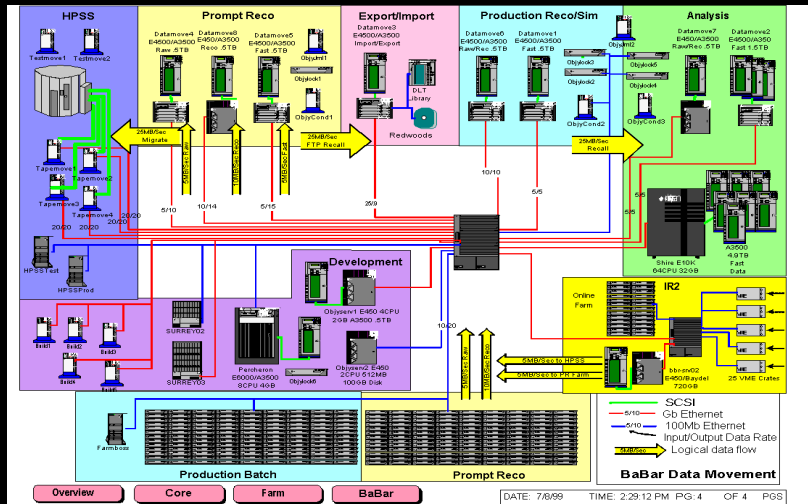


Figure 41

Business Transaction Grid

High-end Transaction Processing Systems

- Example
 - Stock Trading system
 - Many feeds from other systems or exchanges. Head-end based approach to throughput, much like a concentrator
 - Scalable
 - Decentralized Services
 - Plug-in more capacity to cope with spikes in demand
 - QoS guarantees required
 - Real-time or near real-time execution
 - Non-repudiation

ORACLE

Figure 42

Oracle Grid Features

- High Availability
 - Protection from failures, disasters, and human errors
 - 24x7 operation with online maintenance
- End-to-End Grid Security
 - Authentication with SSO/ PKI, Kerberos, and RADIUS
 - Enterprise level authorization and delegation with Enterprise User Security
 - Secure transport via SSL
- Portability between Grid phases

ORACLE

Figure 43

Oracle Grid Features

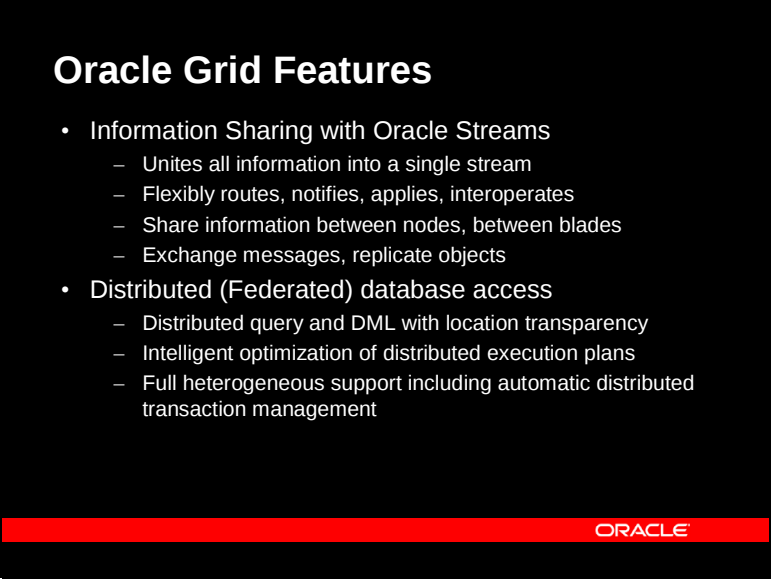
- Manageability
 - Centralized security credential and user management via LDAP
 - Resource Management to enforce fair allocation of database resources
 - Automatic memory SQL execution memory tuning
 - Memory tuning advisors
 - Automatic storage management
 - Enterprise Manager GUI manages complete enterprise stack

ORACLE

Figure 44

ORACLE GRID FEATURES

Transportable tablespaces: Add tablespace to a database and begin processing. Similar to tape racks on IBM mainframes: The database is a tape rack, and the tapes are databases.

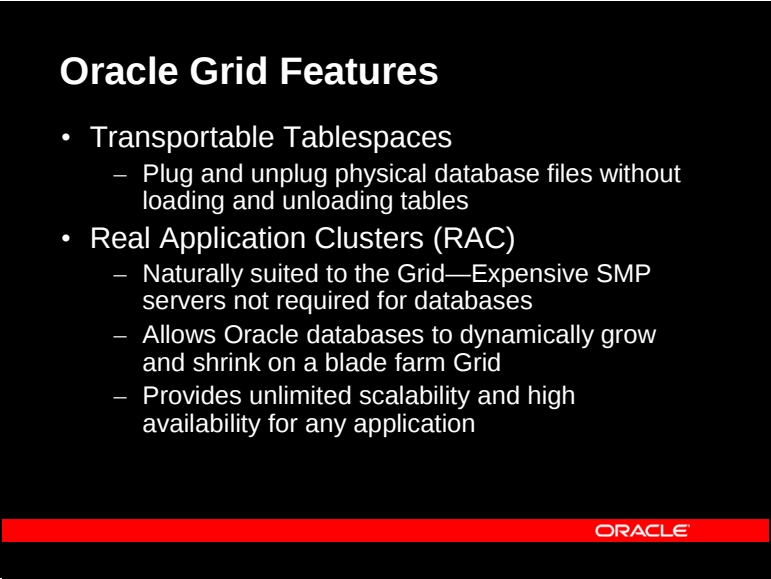
A presentation slide with a black background and white text. The title 'Oracle Grid Features' is at the top. Below it are two bullet points: 'Information Sharing with Oracle Streams' and 'Distributed (Federated) database access', each with sub-bullets. An orange bar at the bottom contains the word 'ORACLE' in white.

Oracle Grid Features

- Information Sharing with Oracle Streams
 - Unites all information into a single stream
 - Flexibly routes, notifies, applies, interoperates
 - Share information between nodes, between blades
 - Exchange messages, replicate objects
- Distributed (Federated) database access
 - Distributed query and DML with location transparency
 - Intelligent optimization of distributed execution plans
 - Full heterogeneous support including automatic distributed transaction management

ORACLE

Figure 45

A presentation slide with a black background and white text. The title 'Oracle Grid Features' is at the top. Below it are two bullet points: 'Transportable Tablespaces' and 'Real Application Clusters (RAC)', each with sub-bullets. An orange bar at the bottom contains the word 'ORACLE' in white.

Oracle Grid Features

- Transportable Tablespaces
 - Plug and unplug physical database files without loading and unloading tables
- Real Application Clusters (RAC)
 - Naturally suited to the Grid—Expensive SMP servers not required for databases
 - Allows Oracle databases to dynamically grow and shrink on a blade farm Grid
 - Provides unlimited scalability and high availability for any application

ORACLE

Figure 46

RAC ARCHITECTURE ADVANTAGE

Early 90's, everyone wrote off shared disk except oracle. Oracle persisted and now owns this space: 45 patents. Competitors can't match. Shared disk matches current trend toward network storage (SAN, NAS) while SN matches limitations of disk storage connectivity from the 80's. Runs real app's : the proof is our customers like UPS (2 x 36 cpu), FAA 5 node Linux, Travelocity, ... and app's like sap, oracle

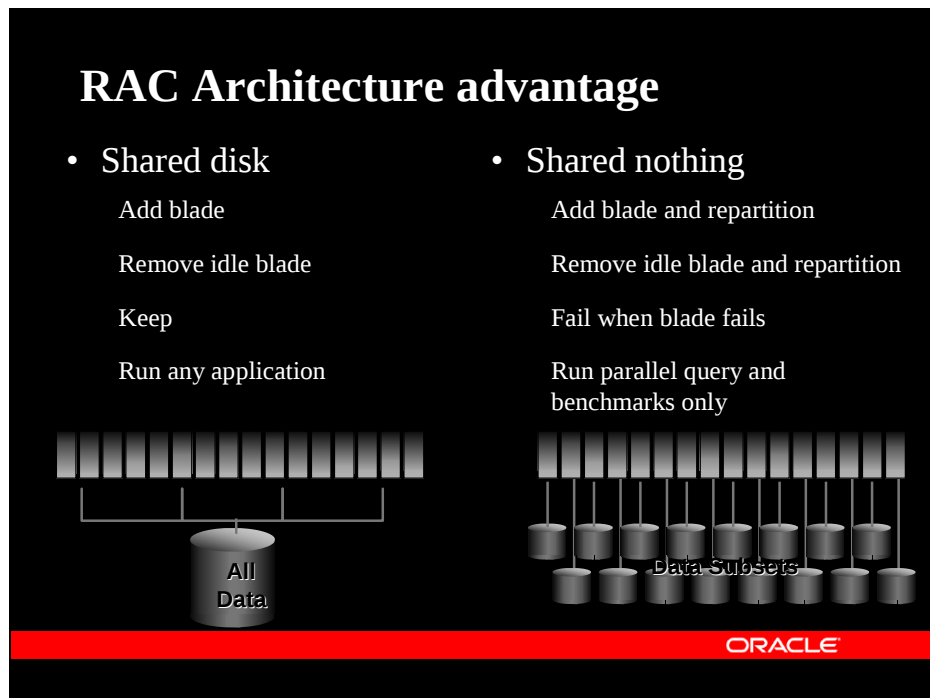


Figure 47

DATASYNAPSE – DO MORE WITH LESS

James Bernardin
DataSynapse, Inc.
New York, NY



Figure 1



Figure 2

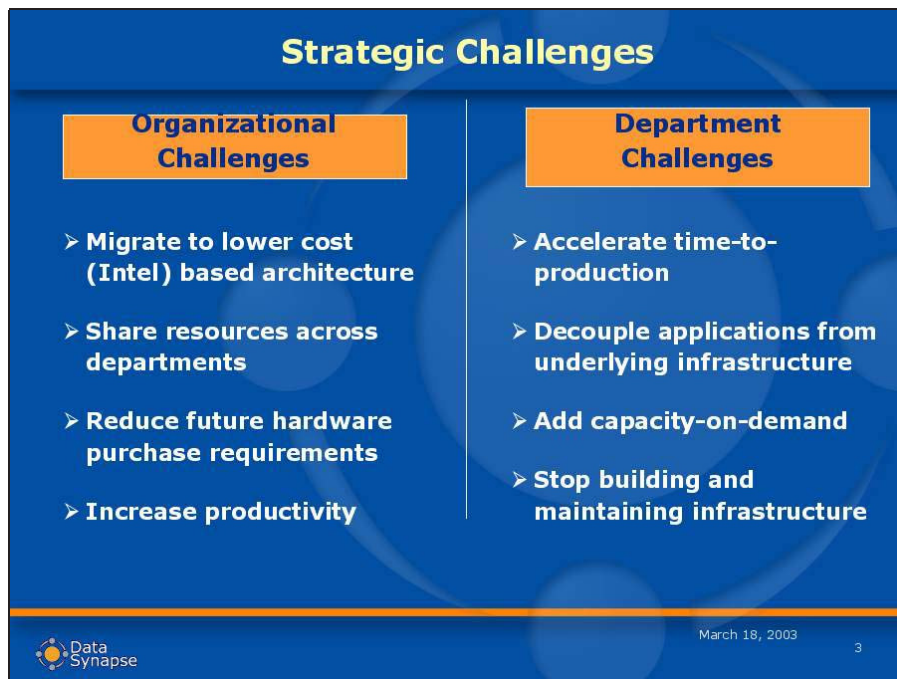


Figure 3

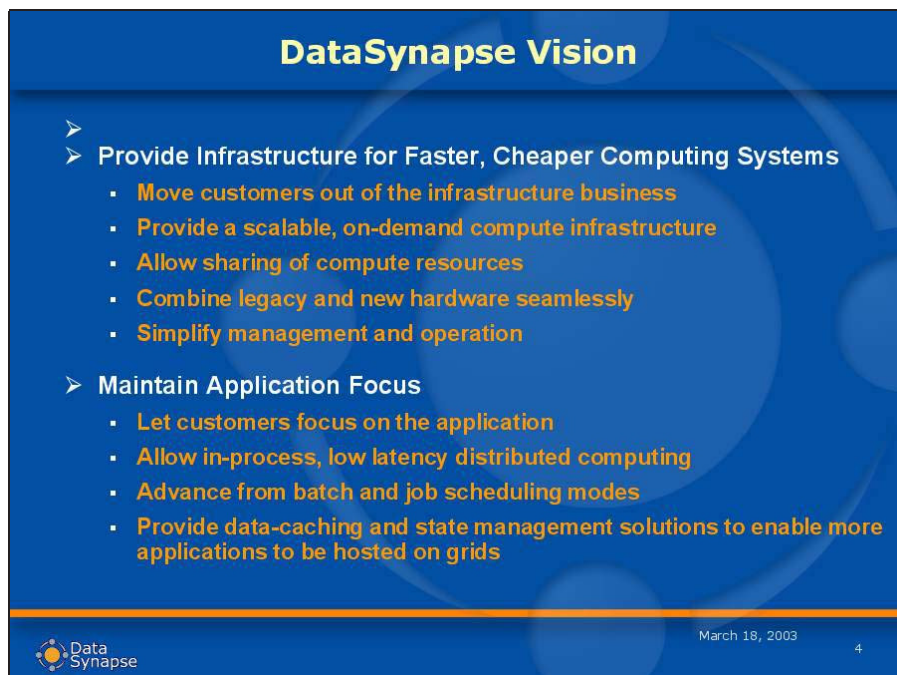


Figure 4

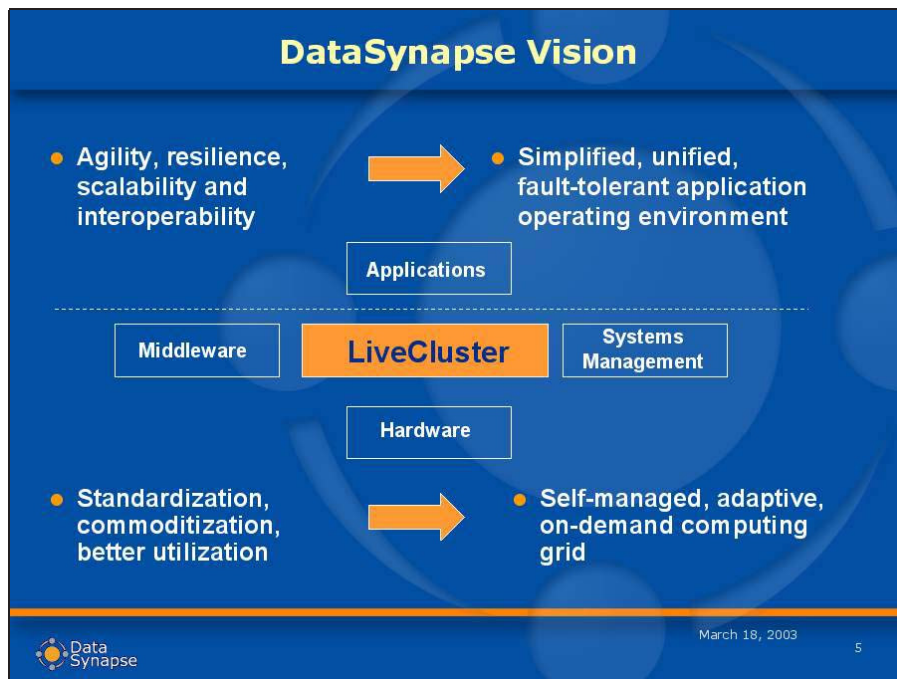


Figure 5

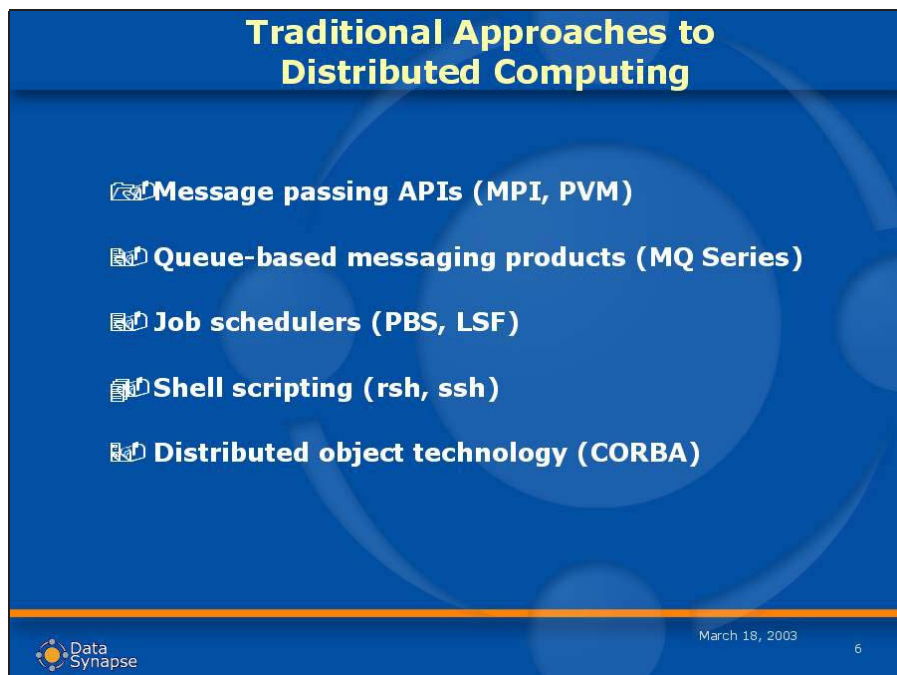


Figure 6



Figure 7

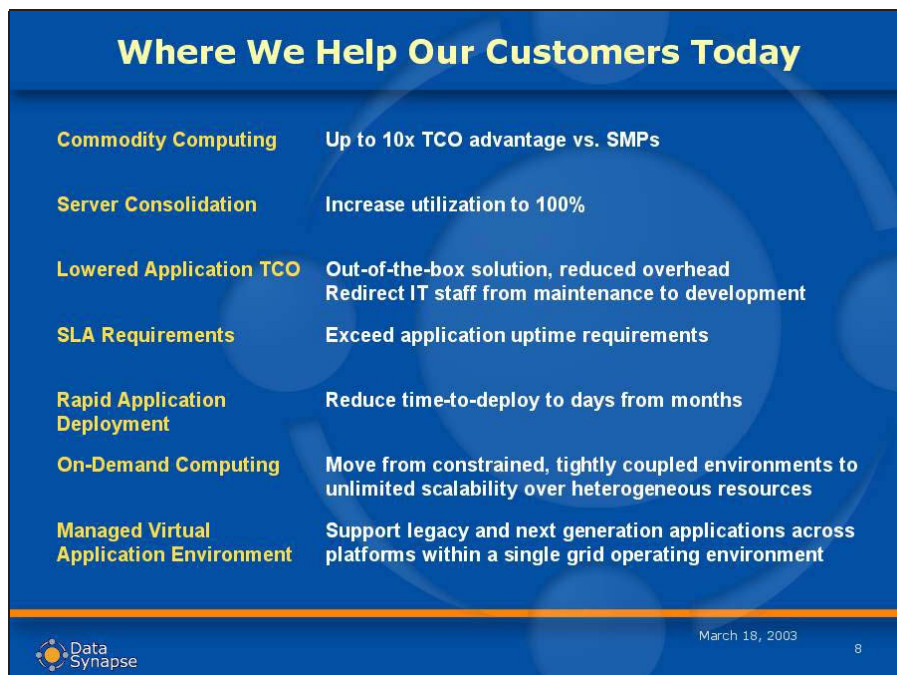


Figure 8

DataSynapse Value Proposition

LiveCluster enables applications requiring scalability to create a virtual environment that transforms IT resources into an on-demand, easily managed grid computing infrastructure

- Dramatically reduce IT cost and application TCO
- Increase application / business performance
- Guarantee application uptime and resilience
- Accelerate time-to-deploy for production systems
- Enable commodity compute models



March 18, 2003

9

Figure 9

Value to Customers

*"Before implementing DataSynapse's LiveCluster solution, running our P&L and risk reports could take as long as 15 hours overnight- **we can now turnaround our mission-critical reports in minutes, on a real-time intraday basis.**"*

*"Moreover, our group is trading 4x more volume and we have increased our modeling simulations by 25x - **about a 100x magnitude performance increase that has scaled effortlessly** on the LiveCluster software platform."*

*"We will trade over \$1 billion of fixed income and related capital markets products over DataSynapse this year- we are booking larger, more exotic, and more lucrative trades with more accurate risk-taking - **DataSynapse helps us make more money, period.**"*

*"**We haven't scratched the surface yet for how we envisage using DataSynapse to meet our ongoing product development and trading activity.**"*

□

Exotics Trader

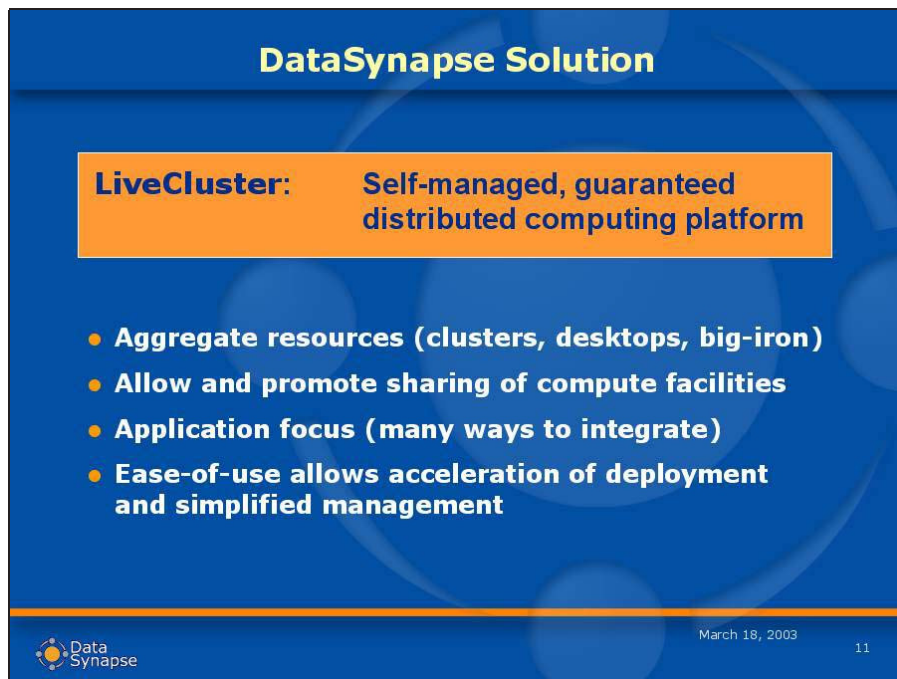
- Andy Cook, Head



March 18, 2003

10

Figure 10

A presentation slide with a blue background and a faint circular pattern. The title 'DataSynapse Solution' is at the top in white. Below it, an orange box contains the text 'LiveCluster: Self-managed, guaranteed distributed computing platform'. A bulleted list follows, and the footer includes the DataSynapse logo, date, and slide number.

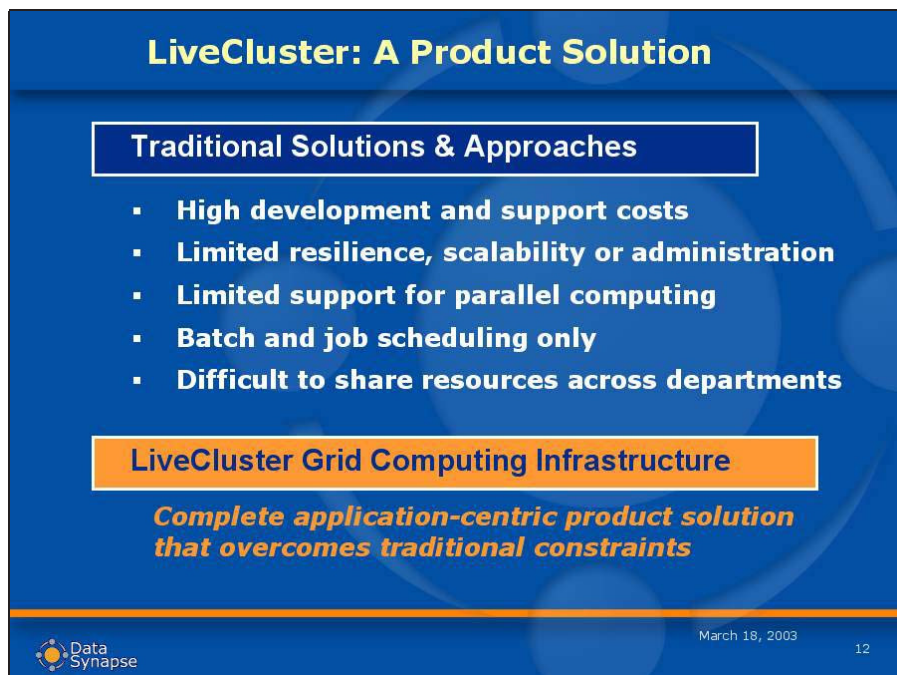
DataSynapse Solution

LiveCluster: Self-managed, guaranteed distributed computing platform

- Aggregate resources (clusters, desktops, big-iron)
- Allow and promote sharing of compute facilities
- Application focus (many ways to integrate)
- Ease-of-use allows acceleration of deployment and simplified management

DataSynapse March 18, 2003 11

Figure 11

A presentation slide with a blue background and a faint circular pattern. The title 'LiveCluster: A Product Solution' is at the top in white. It compares 'Traditional Solutions & Approaches' with 'LiveCluster Grid Computing Infrastructure'. The footer includes the DataSynapse logo, date, and slide number.

LiveCluster: A Product Solution

Traditional Solutions & Approaches

- High development and support costs
- Limited resilience, scalability or administration
- Limited support for parallel computing
- Batch and job scheduling only
- Difficult to share resources across departments

LiveCluster Grid Computing Infrastructure

Complete application-centric product solution that overcomes traditional constraints

DataSynapse March 18, 2003 12

Figure 12

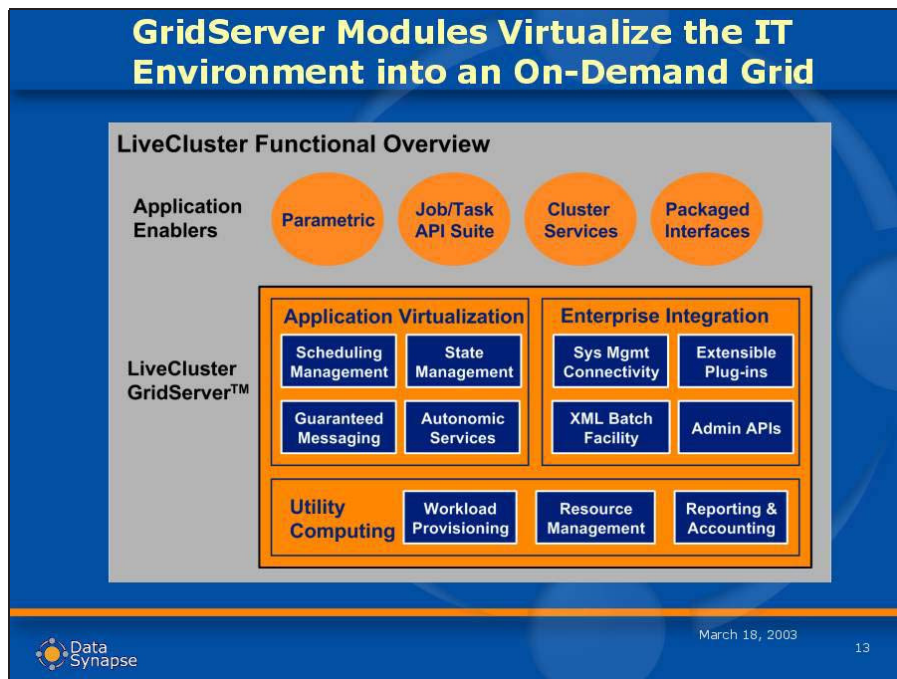


Figure 13

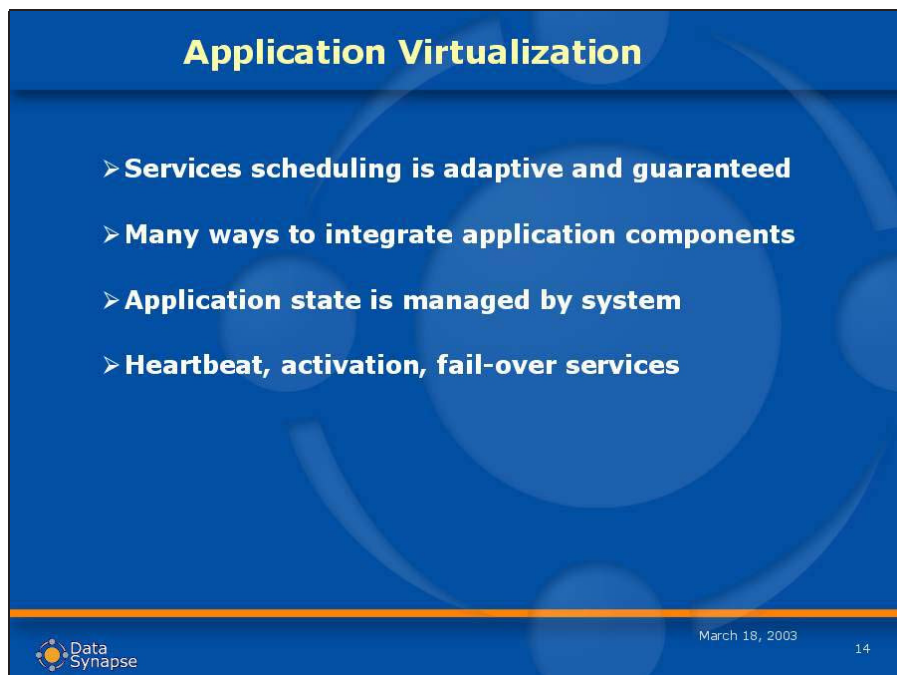


Figure 14

Utility Computing

- Many methods for sharing and provisioning grid resources
- Adaptive scheduling allows for automated scalability
- Utilization statistics and charting (charge-back)
- Audit trail and real-time diagnostics
- Service-based architectures

Data Synapse March 18, 2003 15

Figure 15

Enterprise Integration

- Interface to existing enterprise infrastructure
- Extensible event and command processing
- XML-based workflow/batch facility
- Security plug-ins
- Open APIs for integration with other management systems

Data Synapse March 18, 2003 16

Figure 16

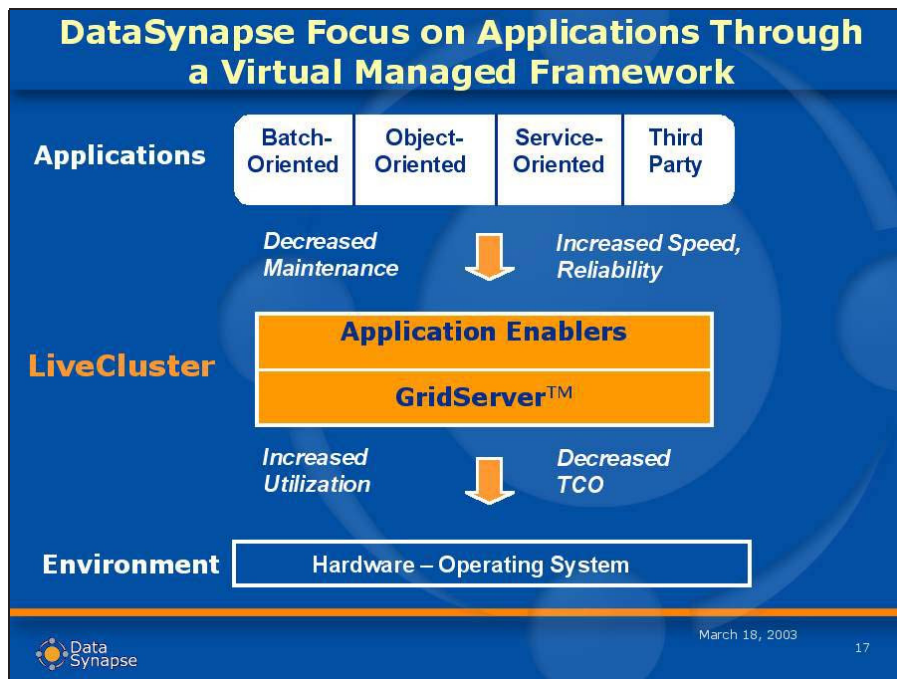


Figure 17

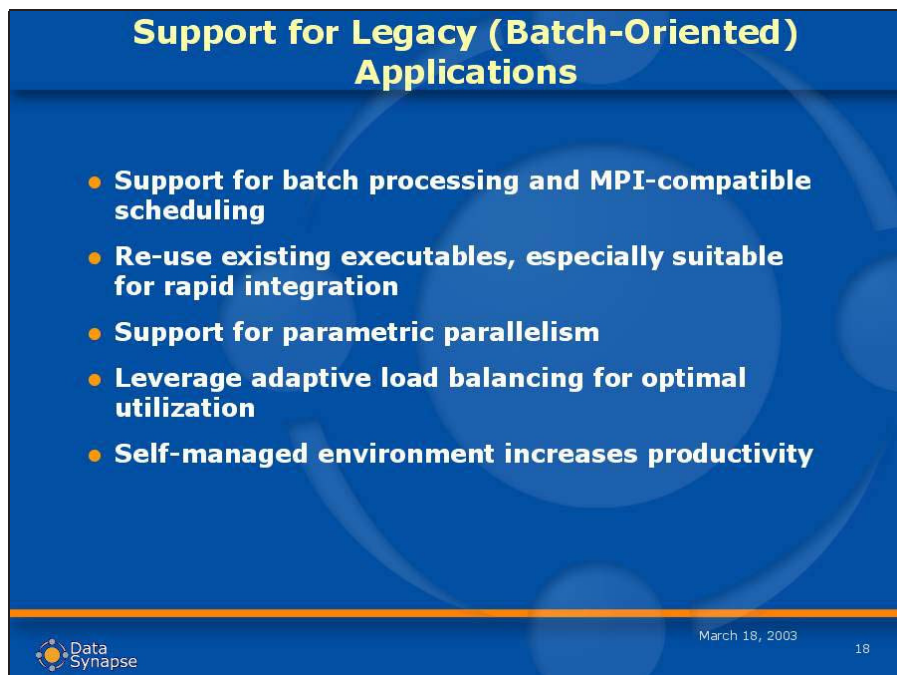


Figure 18

Support for Object-Oriented Applications

- **Support for interactive, GUI-based (e.g. compute- and data-intensive) applications**
- **In-process capability to divide large workloads to perform work in parallel**
- **Based on high-level APIs and clearly defined, distributed object abstractions**
- **Can integrate legacy and next generation applications within days**
- **Improves application performance and guarantees application execution**

 March 18, 2003 19

Figure 19

Support for Service-Oriented Applications

- **Provides resilient, scalable architecture for web services**
- **Suitable for parallel processing and load sharing**
- **Loosely coupled, easy integration method to grid enable application components**
- **Offers explicit support for stateful processing**
- **Uniform cross-language support with simple object-oriented APIs**
- **Requires no language binding with client applications**

 March 18, 2003 20

Figure 20

Corporate Background

- **Management and Foundation:**
 - Founders ex-Wall Street (physicists with NASA heritage)
- **Investors:**
 - Bain Capital, Intel Capital, Wachovia Strategic Ventures
- **Industry Focus:**
 - Finance – Energy - Government
- **Strategic Partners:**
 - Oracle – IBM – Sun – Intel - HP
- **Offices:**
 - New York – London
 - Washington, DC – San Francisco –Houston - Chicago


March 18, 2003
21

Figure 21

Partners Selecting DataSynapse

ISVs:











Infrastructure:







"Grid" Computing Initiatives

- Distributed Computing
- Utility Data Center
- Grid Computing
- N1 Initiative
- Process Area Network


March 18, 2003
22

Figure 22

Intel Selected DataSynapse as its Exclusive Distributed Computing Partner

The Changing **intel.** INFRASTRUCTURE of Business

With today's shaky economic situation and geo-political uncertainties, unpredictable change is one of the few certainties in businesses. How can companies adjust to, and even take advantage of change, rapidly, without disrupting critical operations? Success requires rapid and graceful adaptation—adaptation of business processes along with the applications and information technologies that implement them.

DATASYNAPSE: UNLEASHING THE POWER OF THE GRID

source of sustainable competitive advantage. >>

THE CHANGING INFRASTRUCTURE OF BUSINESS 1

INSIDE

- IBM 2a
- DATASYNAPSE 3a
- BEA SYSTEMS, INC. 4a
- MACROMEDIA 5a
- VERITAS 6a
- BMC SOFTWARE 6a
- BORLAND 7a
- COMPUTER ASSOCIATES 7a

March 18, 2003
23

Figure 23

Industry Analysts Recognize DataSynapse

“By focusing on the applications themselves rather than the traditional approach of focusing on resources, **LiveCluster transforms applications' performance, reliability and resiliency, and cost of ownership.”**

“DataSynapse is leading the way in providing commercial solutions for application reliability and resiliency with its distributed computing solution, LiveCluster.”

“IDC views the DataSynapse product as a **solution for organizations with a variety of IT assets from clusters and hallway grids to campus grids and intraprise grids.”**

“DataSynapse knows how to **easily integrate a guaranteed distributed computing solution into customers' legacy and new applications to exploit an underlying grid environment.”**

“DataSynapse is the **only vendor...that approaches grid computing from a 'commercial' perspective.”**

March 18, 2003
24

Figure 24

DataSynapse and Grid Industry Standards

- LiveCluster is built using accepted industry standards (.NET, Web Services, J2EE)
- Active members of OGSA-WG and is implementing many relevant OGSA services
- LiveCluster is designed to compatible with emerging grid standards

 March 18, 2003 25

Figure 25

Conclusion

- Provide Infrastructure for Faster, Cheaper Computing Systems
 - Move customers out of the infrastructure business
 - Provide a scalable compute infrastructure
 - Allow sharing of compute resources
 - Combine legacy and new hardware seamlessly
 - Simplify management and operation
- Maintain Application Focus
 - Let customers focus on the application
 - Allow in-process, low latency distributed computing
 - Advance from batch and job scheduling modes
 - Provide data-caching and state management solutions to enable more applications to be hosted on grids

 March 18, 2003 26

Figure 26

STAR BRIDGE SYSTEMS

Jim Yardley
Star Bridge Systems
Midvale, Utah

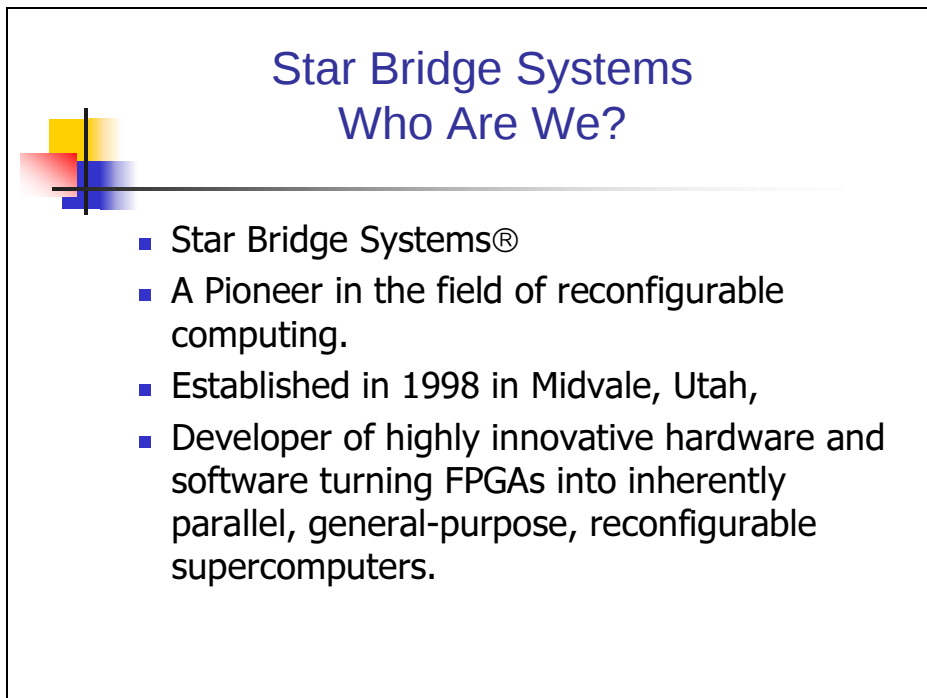


Figure 1

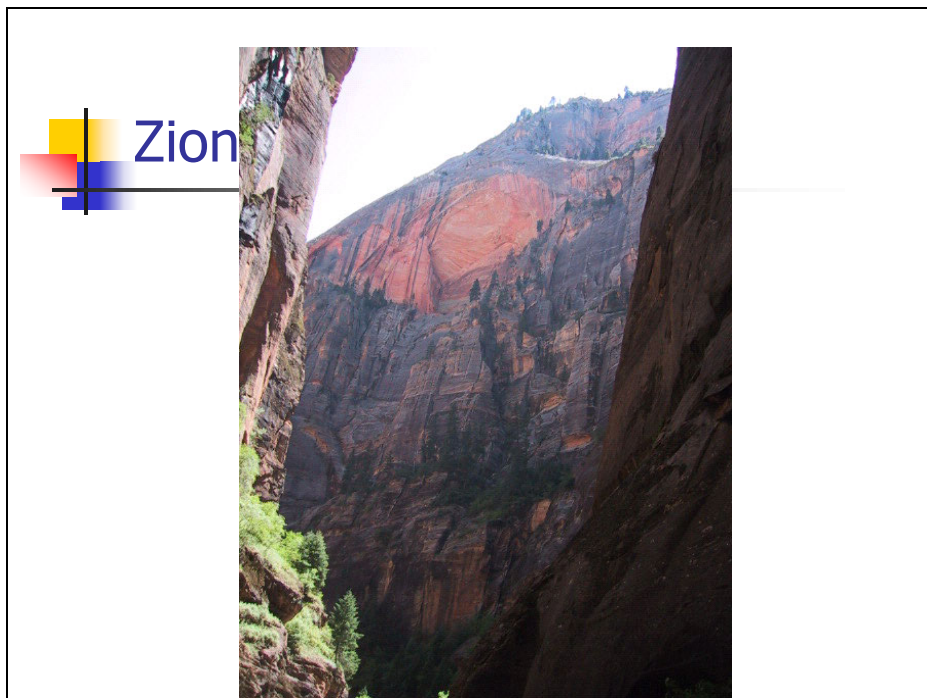


Figure 2

Zion National Park

- Sandstone Monuments towering hundreds of feet over the Virgin River
- Virgin River
 - Head waters 10K feet
 - Zion Park 2000 feet
- Transition of the Virgin River
 - High mountains to Sandstone monuments
 - Narrow canyons through which the water passes

Figure 3

Zion National Park



Figure 4

Zion National Park

- Zion Narrows
 - 17 mile hike
 - Wall to wall water
 - Sandstone cliffs several hundred feet straight up
- Vegetation
 - Trees and bushes growing out of the sandstone walls

Figure 5

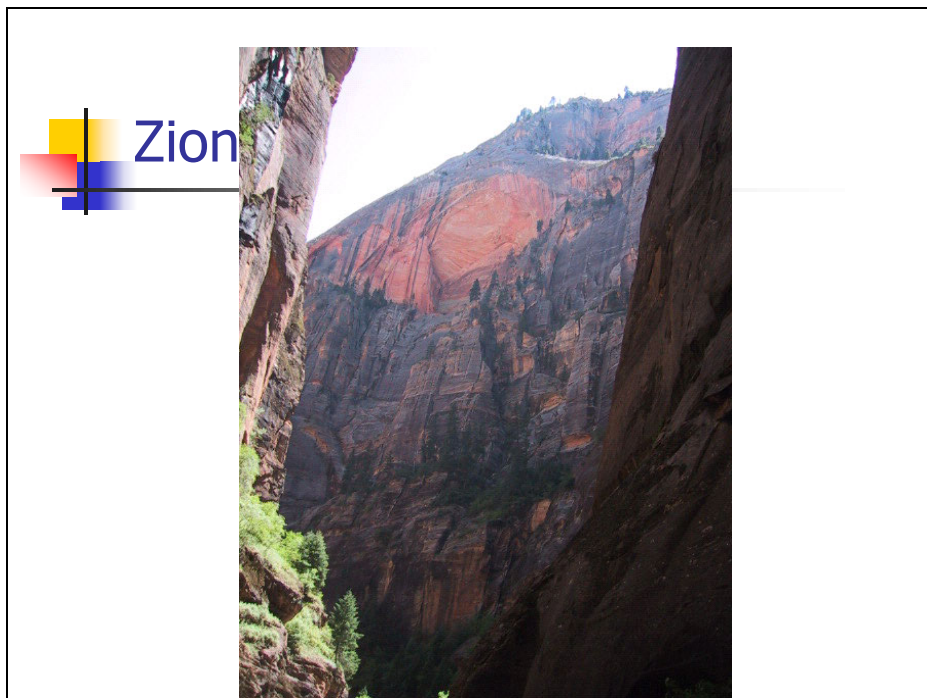


Figure 6



Figure 7

The Tree

- Seed found soil/nourishment
- Roots overgrew their support
- Tree died
- "Dead Tree Syndrome"

Figure 8

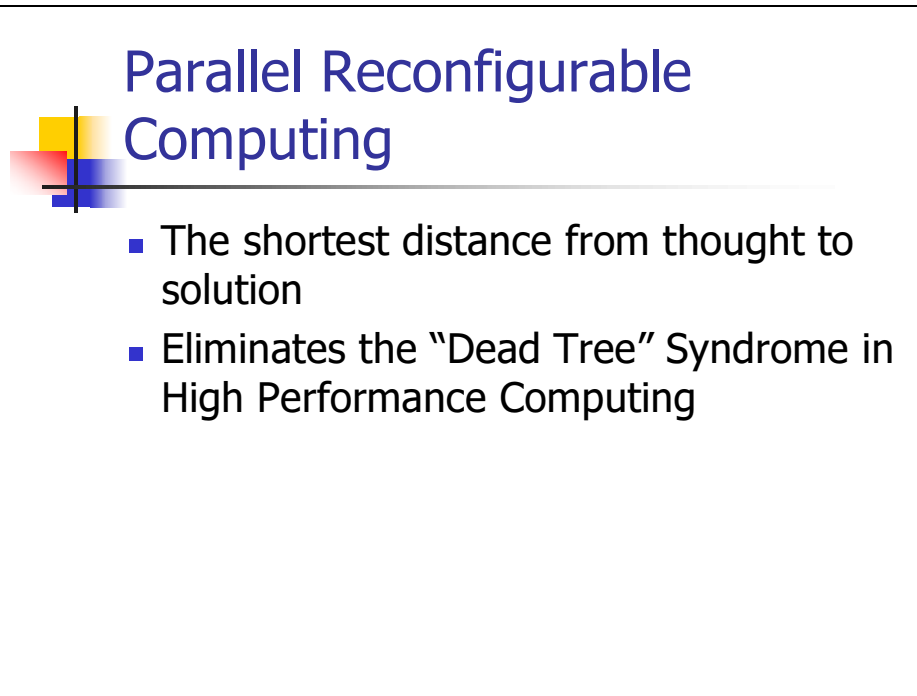


Figure 9

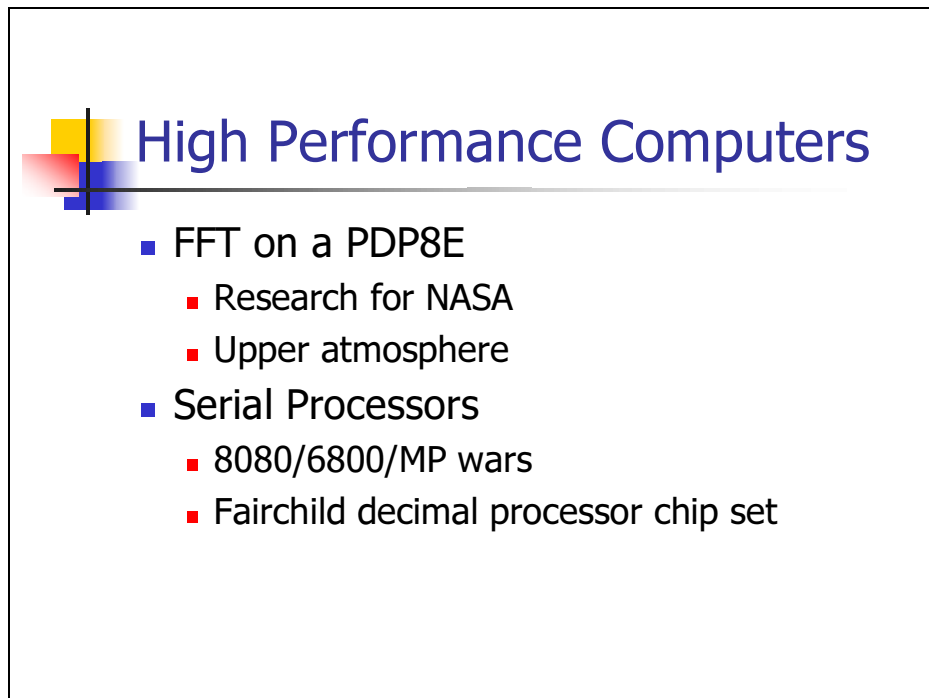


Figure 10

IBM PCs/Clone

- Operating Systems/Databases
 - Microsoft
 - Oracle
- Cluster computers
- Specialized Hardware
- Parallel computing
- COTS—Commercial off the shelf Systems

Figure 11

Clusters and Grids



Figure 12



ASICs

- Cost \$5M to \$40M
- Requires a very large market

Figure 13



FPGAs

- How to program
- Cost to program
- Full data set implementation

Figure 14

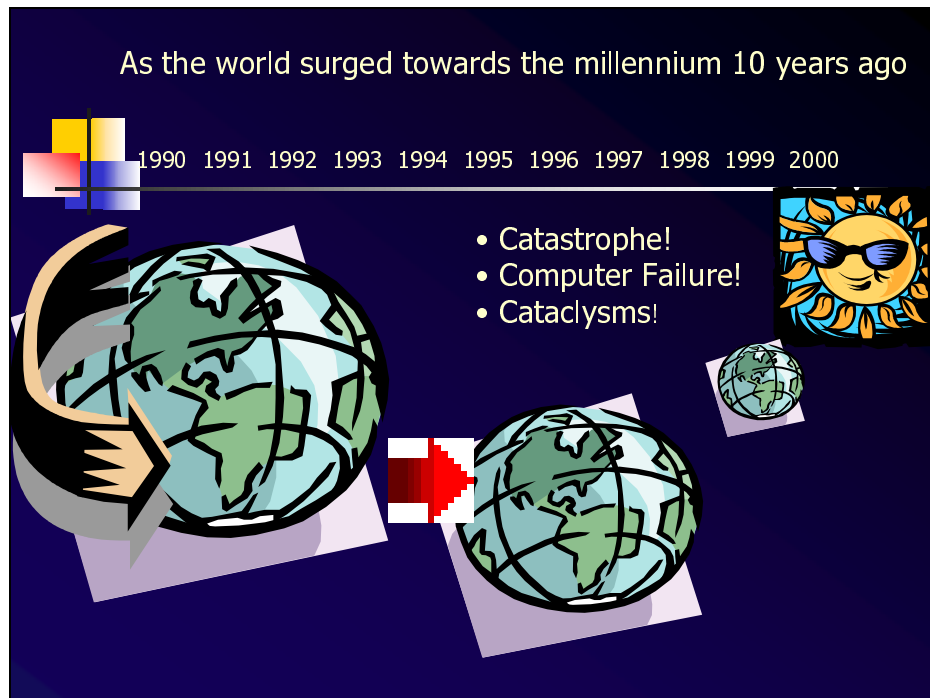


Figure 15

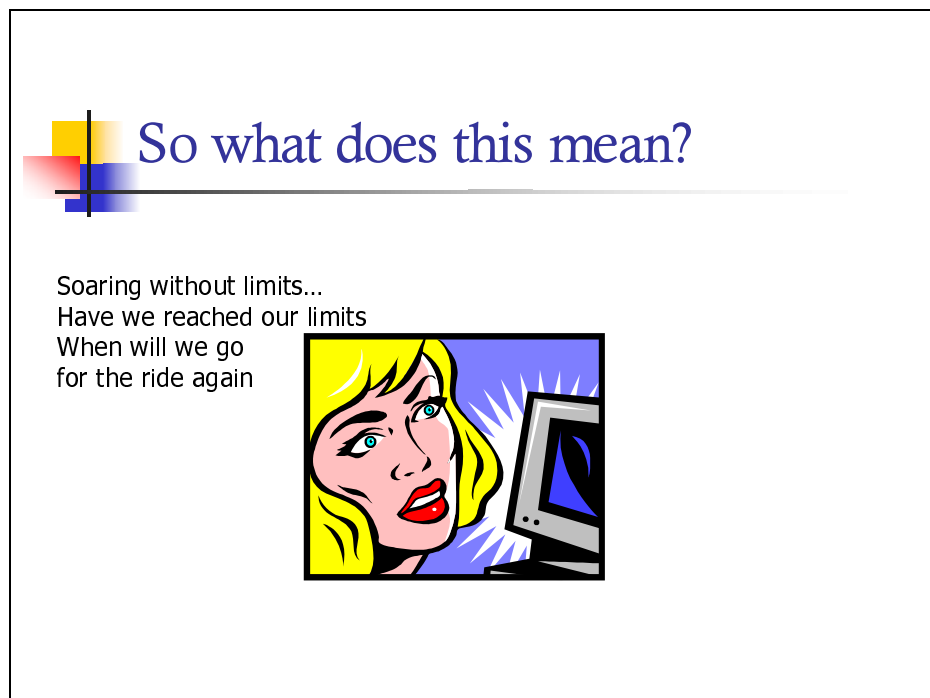
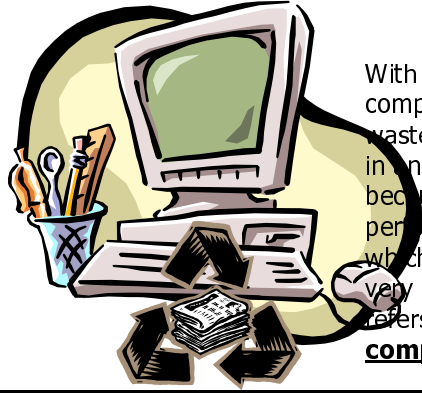


Figure 16

“Re-configurable Computing”,

A phrase coined by Kent Gilson, refers to “the frequent remanufacture or morphing of the entire physical hardware, according to the demands of the user’s specific behavioral requirements”.



With re-configurable computing, you don’t waste a lot of time moving in and out of memory, because all operations are performed on hardware, which makes things move very quickly. Kent Gilson refers to this as **hyper-computing**.

Figure 17

The Potential For Efficient Computing is Greater

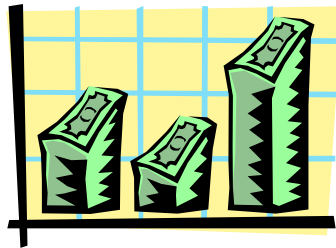
Obtaining parallelism in processing would be a gigantic leap in programming, because it more closely depicts how things happen in the real world.



Figure 18

The Potential For Efficient Computing is Greater

FPGA's are also produced much faster than the standard ASIC chip, (Application Specific Integrated Circuit), which makes them a good choice for the future of hyper-computing. (i.e. it can take a standard ASIC chip up to as long as 18 months to be produced.)



(Decrease production time + Increase in flexibility + Multi-processing) =

Cheaper, Faster, More Efficient Computing

Figure 19

The Need for Faster Machines

- **Problems Dependent on Computation and Manipulation of Large Amounts of Data**
 - Image and Signal Processing
 - Entertainment (Image Rendering)
 - Database and Data Mining
 - Seismic
- **Grand Challenge Problems:**
 - Climate Modeling
 - Fluid Turbulence
 - Pollution Dispersion
 - Human Genome
 - Ocean Circulation
 - Quantum Chromodynamics
 - Semiconductor Modeling
 - Superconductor Modeling
 - Combustion Systems
 - Vision & Cognition

Figure 20



Parallel Computing

- Why? ---- Need for Speed
- What? -----
 - Clusters
 - ASICs
 - FPGAs
 - Heterogeneous
- When? ---- Now
- How? ----- I'm going to show you

Figure 21



Parallel Processing

- Traditional computers- Serial Processing
- HAL hypercomputer- Parallel Processing

Parallel Processing- the ability to execute numerous task simultaneously

- Possible because of FPGAs

Figure 22

This ...



Figure 23

Replaces This !



Figure 24



Viva, the OS of Reconfigurable Computing, is a graphical programming interface that uses an object oriented programming model with data flow characteristics.

Reusable objects can be built up from primitive functions or multiple levels of hierarchy. Programs are written and compiled in Viva as well.

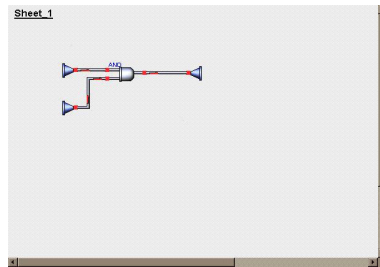
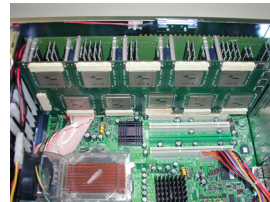
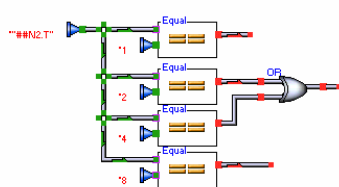


Figure 25



Star Bridge Systems Fundamental Technology Shift

*"This technology will allow us to translate our
ideas into solutions ... as fast as we can think"*
Dr. Robert Singletary, NASA Research Scientist



High Level Language

- High level graphical language
- Directly programs FPGAs
- Reusable highly flexible objects

True Parallel Hardware

- Reconfigurable hardware
- Unique FPGA Implementation
- Perfect computing every time

Figure 26

Speed Through Superspecificity

- Microprocessors are usually designed with generalized functions to address a wide variety of algorithmic applications.
- Star Bridge technology creates only the necessary and sufficient circuitry needed for the specific application.
- FPGA core development – 1/16th to 1/20th the size of VHDL designed cores.

Figure 27

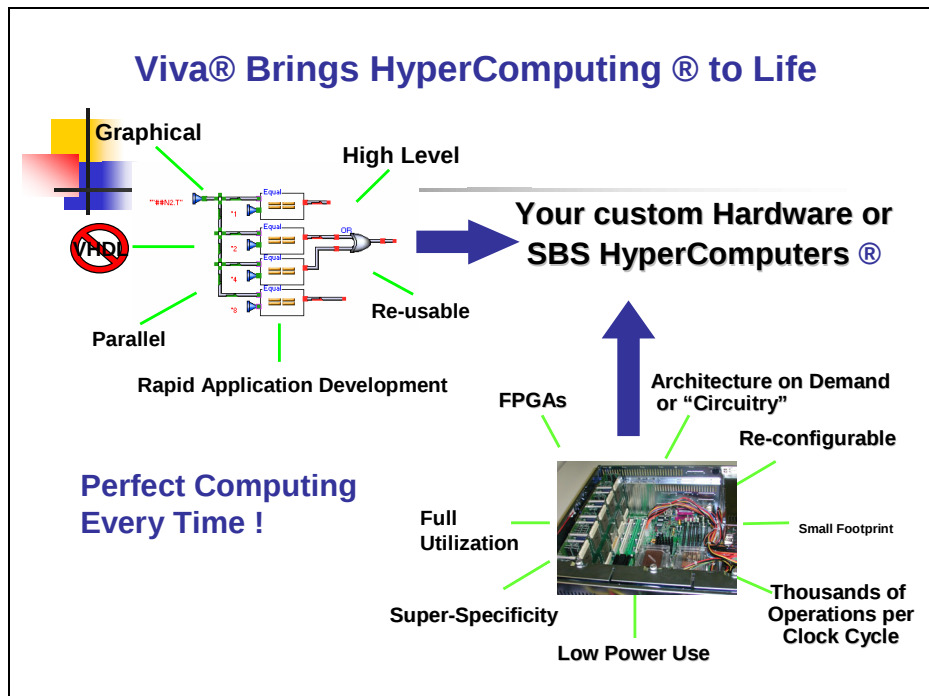


Figure 28

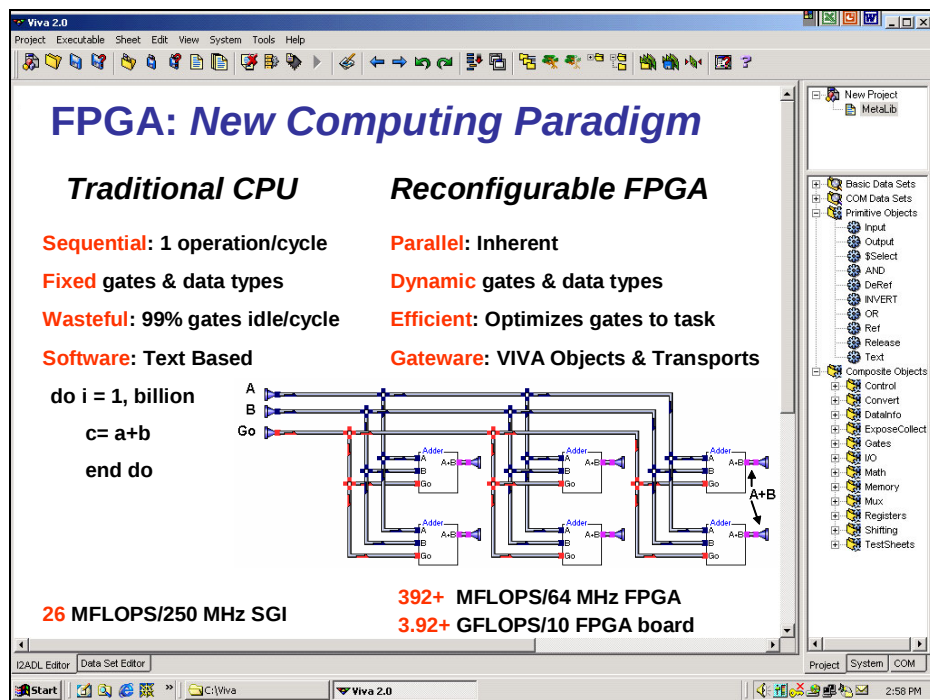


Figure 29

“Our mandate is to pioneer the future ... to push the envelope ... to do what has never been done before.”

NASA Vision by Administrator Sean O’Keefe, April 12, 2002 , [Maxwell School at Syracuse University](#)

- ILLIAC (Ames-1972)
- Finite Element Machine (Langley – 1982)
- MPP (Goddard-1983)
- Cray (Langley-1989)
- Intel (Delta – 1992)
- IBM (LaRC – 1995)
- Star Bridge Systems Hal15 (Langley – 2002)
- Star Bridge Systems HC36 (Langley– 2003)

Figure 30



Traditional Sequential Programming

- Fortran and C programs have been written for serial computers:
 - One instruction executed at a time
 - Using one processor or clusters
 - Processing speed dependent on how fast data can serially move through hardware and subsystems
 - Subsystem communications bottlenecks

Figure 31



How to Program Reconfigurable Computers?

- Data Flow Programming Style
- Must Think "Inherently Parallel"
- Graphical Based "Language"
- Everything Tied to Clock Signals
- No Von Neumann Bottlenecks
- Programming Power Tied to Number of Gates or Area (Number) of FPGAs

Figure 32



Hal-15

- Present Product
 - Hal-15
 - 10 FPGAs/Board
 - 20 billion MAC(16-bit multiply accumulates/sec)
 - 5 billion FLOPS (32-bit Floating Point Operations/sec)
 - 500 Giga OPSS (4-bit, integer operations/sec)
 - Configurable Options:
 - 2-10 FPGAs per Board, up to 20 boards per system
 - Viva 1.5 Release

Figure 33



The 3-Points of HAL: Hyper Algorithmic Logic Computer

- The HAL 15 system here at NASA was the first system to be delivered by Star Bridge to an established high performance computer user. The HAL 15 uses a combination of an Intel based workstation, and a PCI board containing 10 Xilinx FPGA chips.
- IIADL (Implementation Independent Algorithm Description Language) a new programming language that makes it possible for an FPGA-based re-configurable computer to operate as a general-purpose computer system.
- Viva (Latin word for "life"), brings life to HAL and hyper-computing as an OS, compiler, and graphical user interface all in one.

Figure 34



HAL in NASA

- Spacecraft and Satellite control centers
- Solutions for structural, electromagnetic and fluid analysis
- Radiation analysis for astronaut safety
- Atmospheric science analysis
- Digital signal processing
- Pattern recognition
- Acoustic analysis

Figure 35



Langley Algorithms Developed*

- **Matrix Algebra:** Vectors, Matrices, *Dot Product*
 - **Factorial** => Probability: Combinations/Permutations **AIRSC**
 - **Cordic** => Transcendentals: sin, log, exp, cosh...
- **Integration & Differentiation** (numeric)
- **Matrix Equation Solver:** $[A]\{x\} = \{b\}$ via Gauss & Jacobi
 - **Dynamic Analysis:** $[M]\{\ddot{u}\} + [C]\{\dot{u}\} + [K]\{u\} + \text{NLT} = \{P(t)\}$
 - **Analog Computing:** digital implementation
- **Nonlinear Analysis:** “Analog” simulation avoids **NLT** solution development time

* In *AIAA & Military & Aerospace Programmable Logic Device (MAPLD) papers*

Figure 36

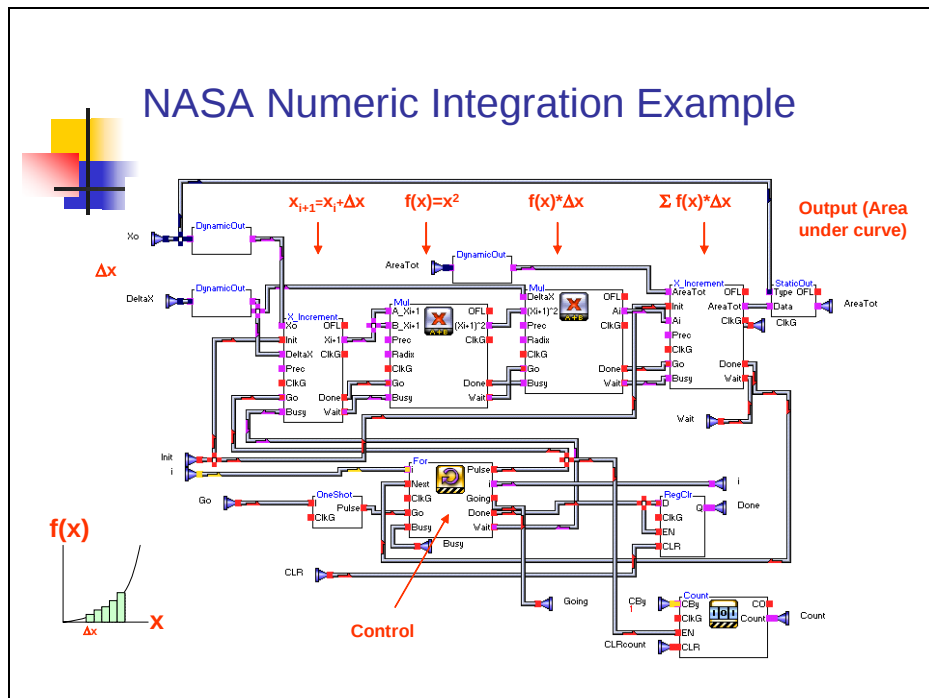


Figure 37

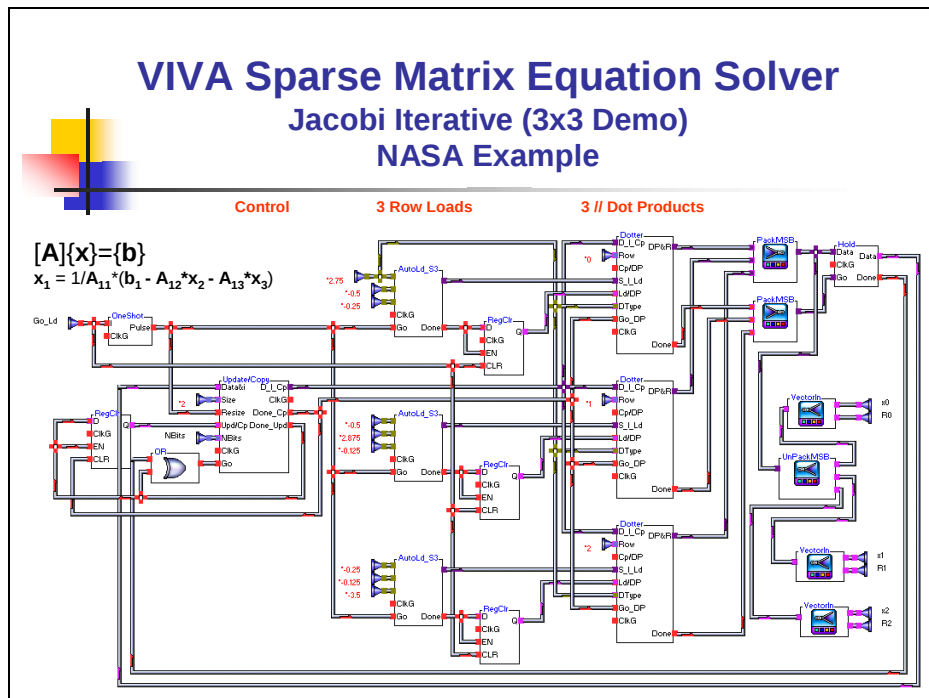


Figure 38

VIVA: Gateware Development Tool

What: Simple tool to configure FPGAs (VHDL cumbersome)

How: Transforms high-level graphical code to logic circuitry

Why: Achieve near-ASIC speed (w/o chip design)

Growth in VIVA Capability

VIVA 1 (Feb '01)	VIVA2 (July '02)
NO Floating Point	Extensive Data Types
NO Scientific Functions	Trig, Logs, Transcendentals
NO File Input/Output	File Input/Output
NO Vector-Matrix Support	Vector-Matrix Support
Access to One FPGA	Access to Multiple FPGAs
Primitive Documentation	Extensive Documentation
Weekly Changes	Stable Development
Frequent "bugs"	Few "bugs"

Figure 39

Progress - Roadmap

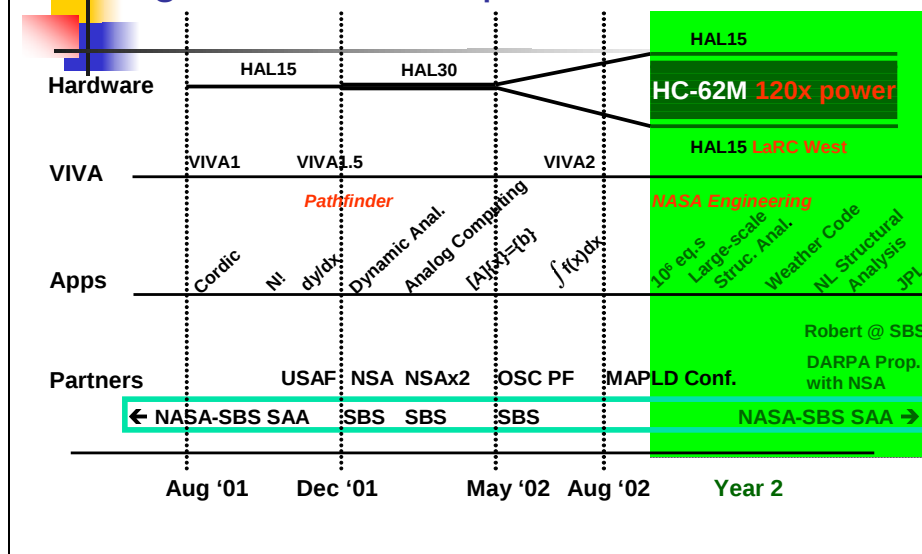


Figure 40

Year 2: Exploit Latest FPGAs

Rapid Growth in FPGA Capability

	FPGA (Feb '01)	FPGA (Aug '02)
Xilinx FPGA	XC4062	XC2V6000
Gates	62K	6 million (97x)
Multiplies in H/W	0	144 (18x18)
Clock Speed MHz	100	300 (3x)
Memory	20Kb	3.5 Mb (175x)
Memory Speed	466 Gb/s	5 Tb/s (11x)
Reconfigure Time	100ms	40ms (2.5x)
GFLOPS	0.4	47 (120x)
Total GFLOPs	4 (10 FPGAs)	470 (10 FPGAs)

Plans:

- Millions of Matrix Equations for Structures, Electromagnetics & Acoustics
- Rapid Static & Dynamic Structural Analyses
- Cray Vector Computations in Weather Code (VT PhD)
- Robert on Administrator's Fellowship at Star Bridge Systems
- Joint proposals with NSA & DARPA
- Simulate advanced computing concepts using VIVA
- Collaborate with SBS to expand VIVA libraries
- Influence VIVA development to meet NASA application needs
- Expand FPGA applications for NASA programs

Figure 41

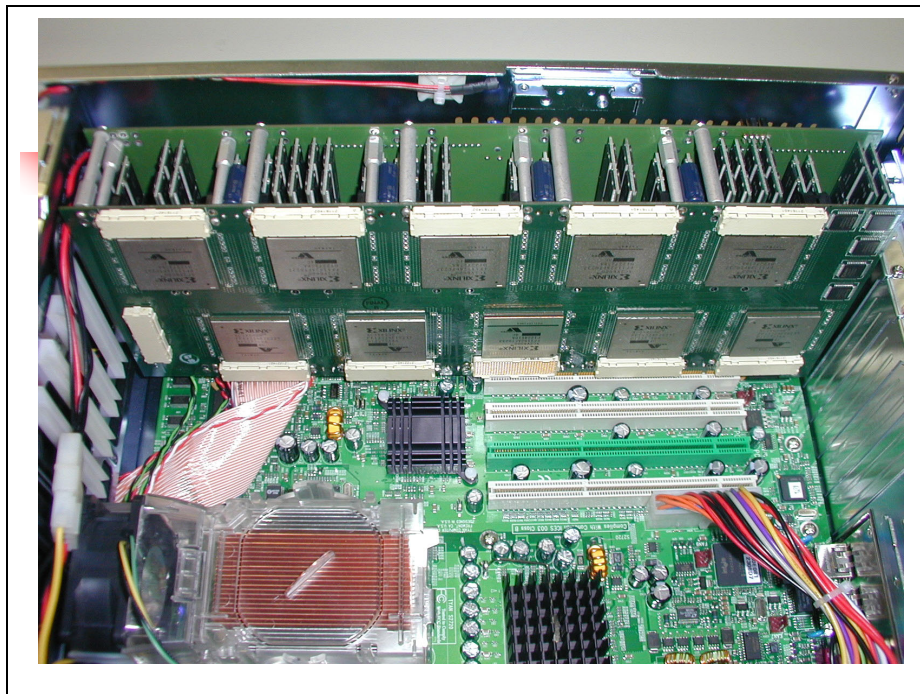


Figure 42



FPGA Performance Curve

- Microprocessor density is doubling only every 18 to 24 months under Moore's Law
- FPGAs are on a much steeper growth curve.
- In late 1999, Xilinx projected the following growth path for its most powerful FPGA chips:
 - 1998 1 million gates per FPGA
 - 1999 2 million gates per FPGA
 - 2000 4 million gates per FPGA
 - 2002 10 million gates per FPGA
 - 2004 50 million gates per FPGA

Figure 43

Virtex-II - Distributed DSP Resources

■ LUTs* & Registers

- Up to 122,880
- Logic + storage

Usage examples

- Pipelined algorithms
- Multiple channels
- Coefficient storage
- Shift registers/delay

■ 18x18 Multipliers

- Up to 192
- 200+ MHz

Usage examples

- High-performance FFT
- Equalizers

■ Active Interconnect™

- Connects all resources

■ Block RAM

- Up to 3.5 Mbit, true dual-port

Usage examples

- Data buffering + storage
- Single-chip 2-8K FFT
- Video line buffers

■ SelectI/O™

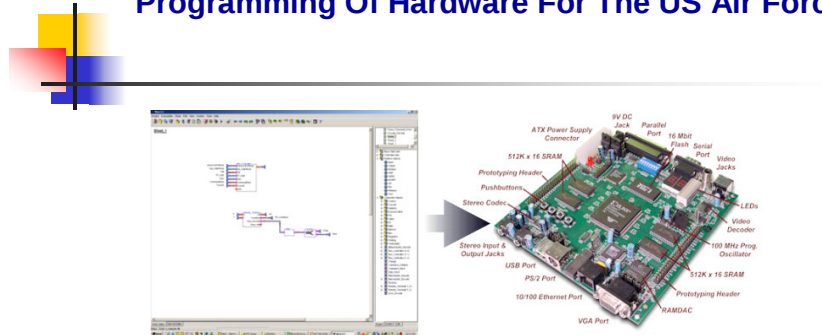
- Up to 1,108 programmable I/Os

* LUT = Lookup Table

www.xilinx.com

Figure 44

Reconfigurable Computing Software Enables Easy Programming Of Hardware For The US Air Force



VIVA Software to Hardware

"Traditional approaches to programming algorithms in FPGA are laborious and time consuming. Star Bridge Systems® of Midvale, Utah developed an Electronics Design Automation (EDA) tool called Viva® that solves these problems"

Mr. Lloyd Reshard, AFRL/MNAV

Figure 45

Star Bridge Systems

- FPGA Technology
 - Inherently Parallel Compute Substrate
- Reconfigurable Hardware
 - How do you easily Program? Or Reconfigure???
- Viva Software
- Hypercomputer
 - Parallel Supercomputer Capacities and Capabilities

Figure 46



Star Bridge Systems

- Reconfigurable Computing Technology
 - Massively and inherently parallel
 - Asymmetrical
 - Ultra-tightly coupled
 - Linearly scalable
- Realization of Parallel Computing

Figure 47



Star Bridge Systems, Inc.

- Hardware and Software Integrated System
 - Parallel and Reconfigurable Computing
 - Greater Programming Flexibility
 - Lower Cost
 - Lower Power Consumption
 - Smaller Size
 - Breakthrough in Performance and Speed

Figure 48



List of Customers

- NSA
- George Washington University
- University of South Carolina
- George Mason University
- North Carolina A&T
- NASA
 - Langley
 - Marshall Space Flight Center
- San Diego Supercomputer Center
- U.S. Air Force Eglin AFB
- National Cancer Institute—Bio-informatics
- Commercial Customers for EDA applications

Figure 49



What are People Doing With This?

- NASA
- Seismic Petroleum Exploration
- Smiths-Aerospace
- National Cancer Institute
- US Air Force

Figure 50

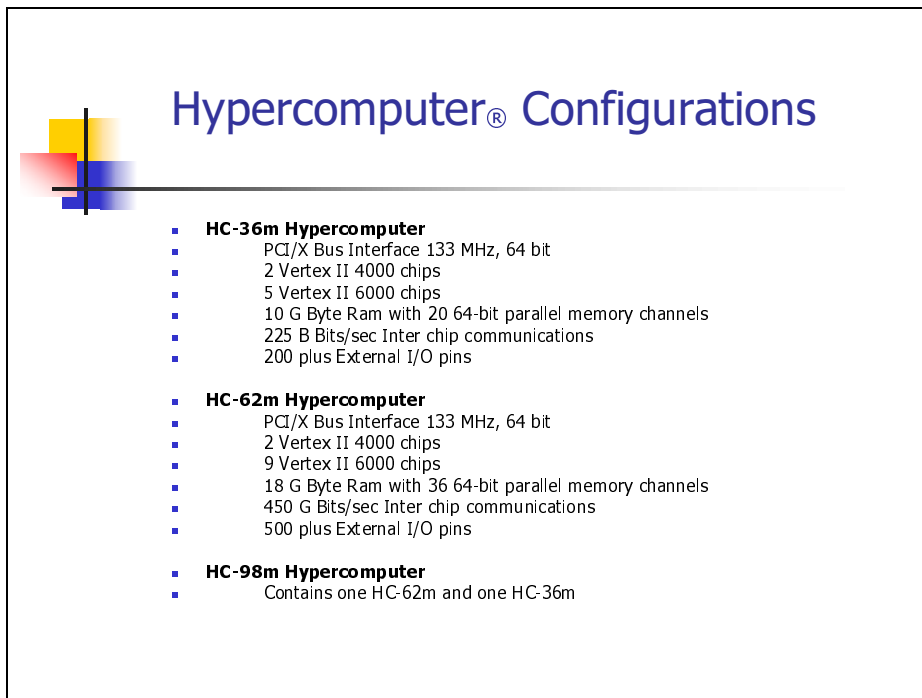


Figure 51

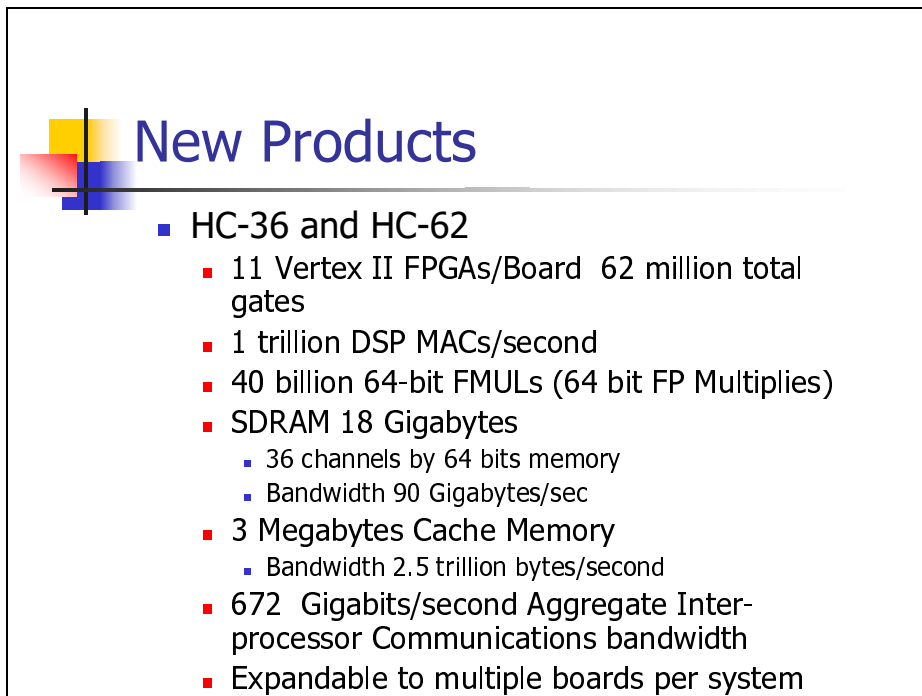


Figure 52



10 Chip Co-Processor Board

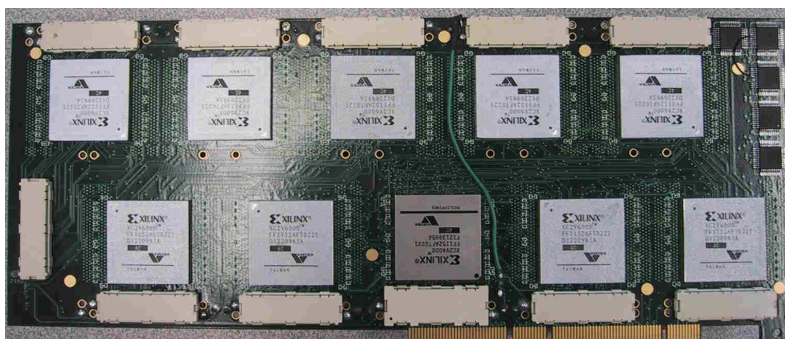


Figure 53



HC-62 Board Set Assembly

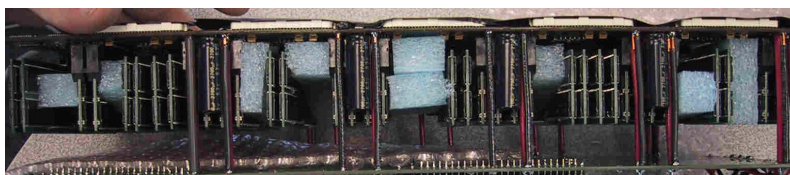


Figure 54



HC-62 Assembled in Dual Xeon Motherboard

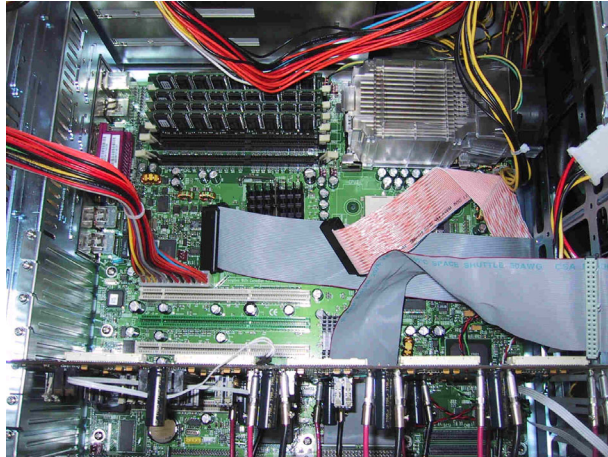


Figure 55



Star Bridge Systems

- Technical Discussion
 - FPGA Hardware configuration
- Viva Discussion
 - Viva Demonstration
- Bio-Informatics System Application
 - Smith Waterman
- Air Force 1553 Standard
- Programming Examples

Figure 56

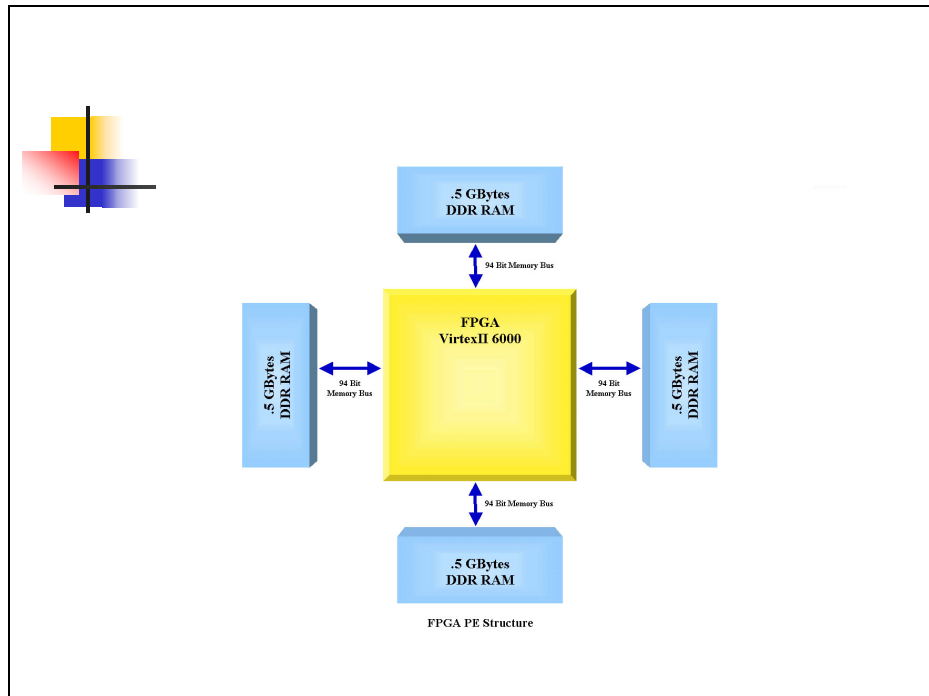


Figure 57

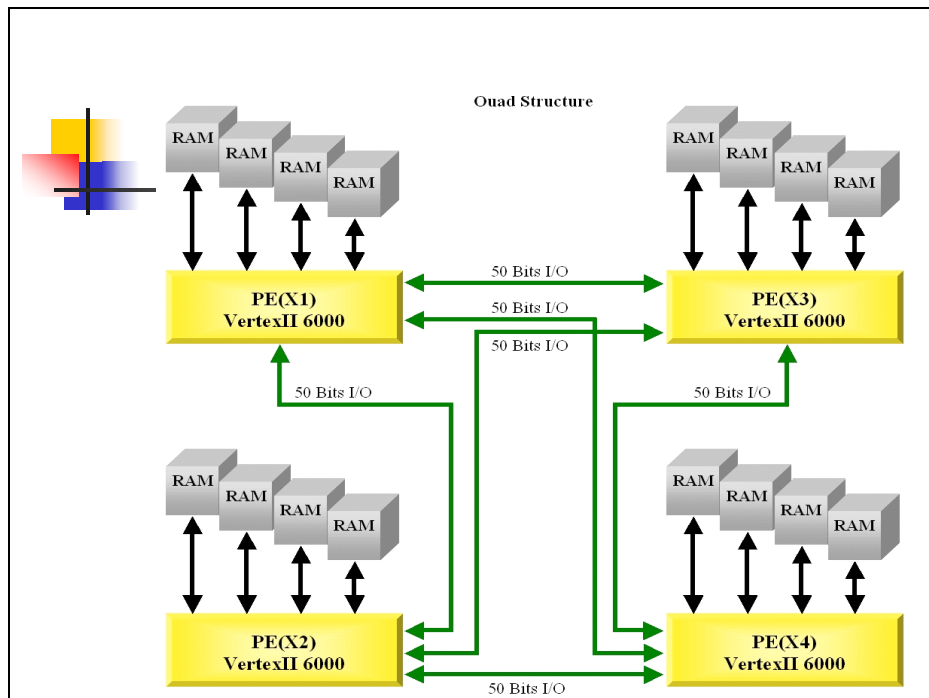


Figure 58

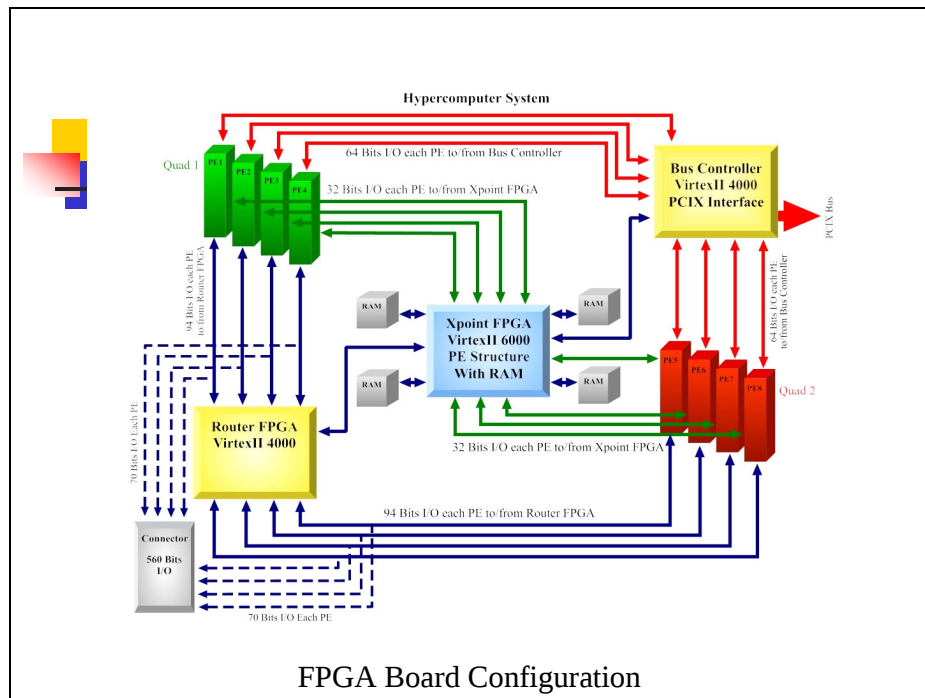


Figure 59



Star Bridge Systems, Inc.

- Viva --- Capability Computing Software
 - Programming Language
 - Compiler/Operating System
 - Graphical User Interface
 - Integrates Hardware and Software
 - Software Implementation at Hardware Speeds
 - Increased Productivity of Application Developers
 - Solves Difficultly in FPGA Program Development

Figure 60

VIVA software

- What: Graphical Programming Language
- How: Transforms high-level graphical code to logical circuitry
- Why: Achieves near ASIC speed



Figure 61

Viva— Primary Elements

- Rapid Application Development Environment
- Parallel Component Object Oriented Language
 - IIADL—Implementation Independent Algorithm Description Language
- Execution Target/System Definition Tools
- Multi-Process Execution and Reconfiguration OS-Kernel
- Application Builder Libraries METALIB
- System Target Libraries
- User Interface/STDIO COM/ActiveX Component Library

Figure 62

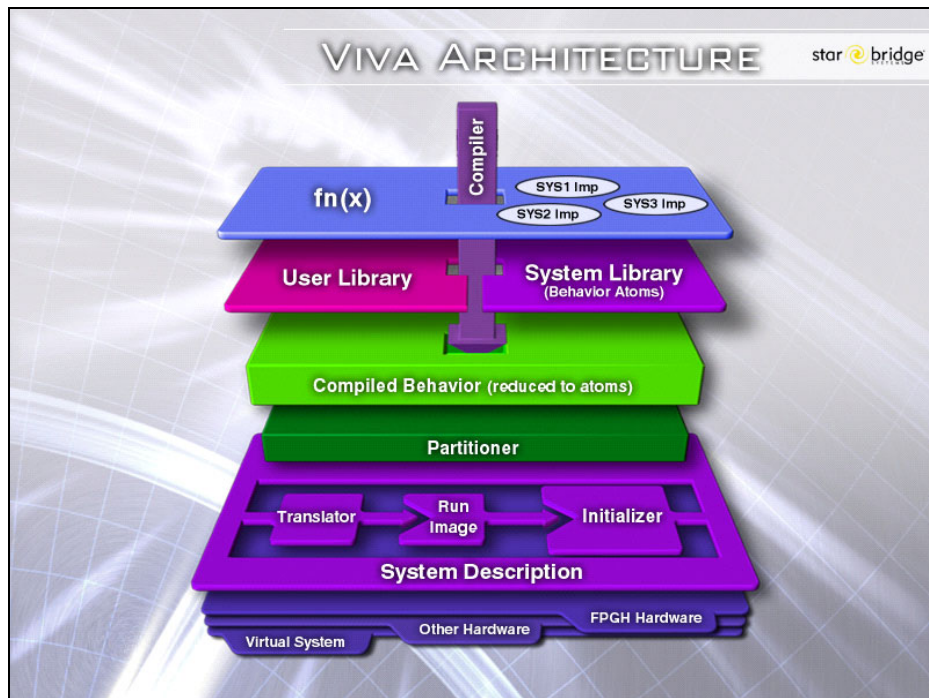


Figure 63

Rapid Application Development

- COM/ActiveX Form Designer
- Drag n Drop Graphical Program Editor
- Drag n Drop Data Set Editor
- Drag n Drop System Builder and Resource Editor
- Auto Generated Widget Interface

Figure 64

IIADL and Compiler

- Data Flow Centric Programming Model
- Parallel Component Object Oriented Language
- Recursive Algorithm/Topology
- Unlimited Operator Overloading
- Data Set Polymorphism
 - Multi-Precise and Multi-type Operators
- Information Rate Polymorphism
 - Multi-Rate Operators
- Context Sensitive Operator Synthesis
- Strong Types
- Data Set Composition/Decomposition Operators
- Dynamic Data Set Creation
- Timing Driven Partitioning/Co-synthesis

Figure 65

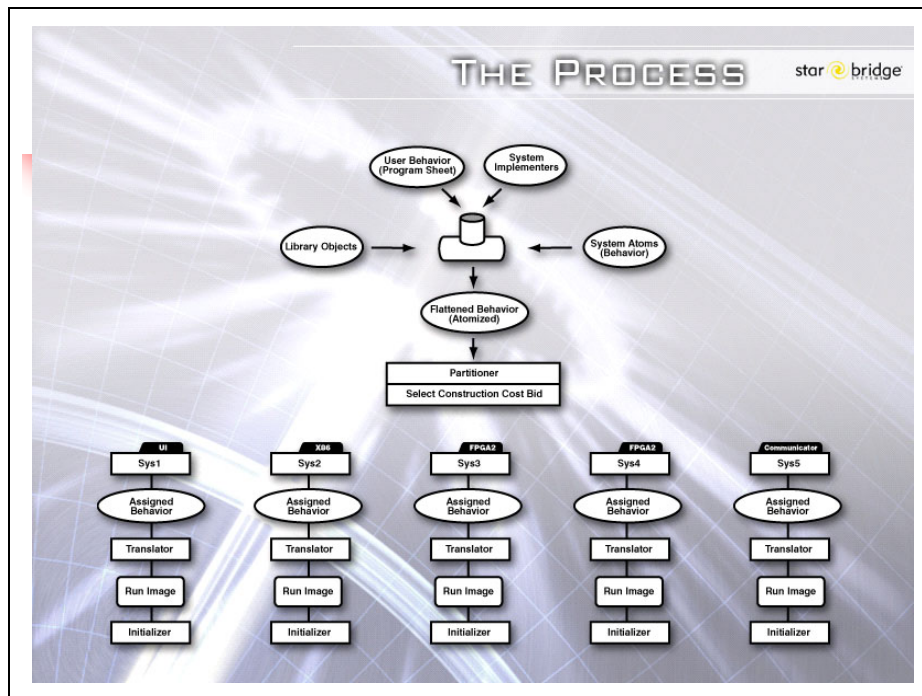


Figure 66

Operating System Kernel

- Event Driven Real-time Mixed-Mode Execution Kernel
- Multi-process Hierarchical Thread creation and lifetime management operators
- Dynamic default Interface Creation and Execution
- Inter-process communication and processor side memory management
- Full COM(Common Object Model) Execution
- Stand Alone Executable Creation for Application Distribution
- Dynamic Reconfiguration Support

Figure 67

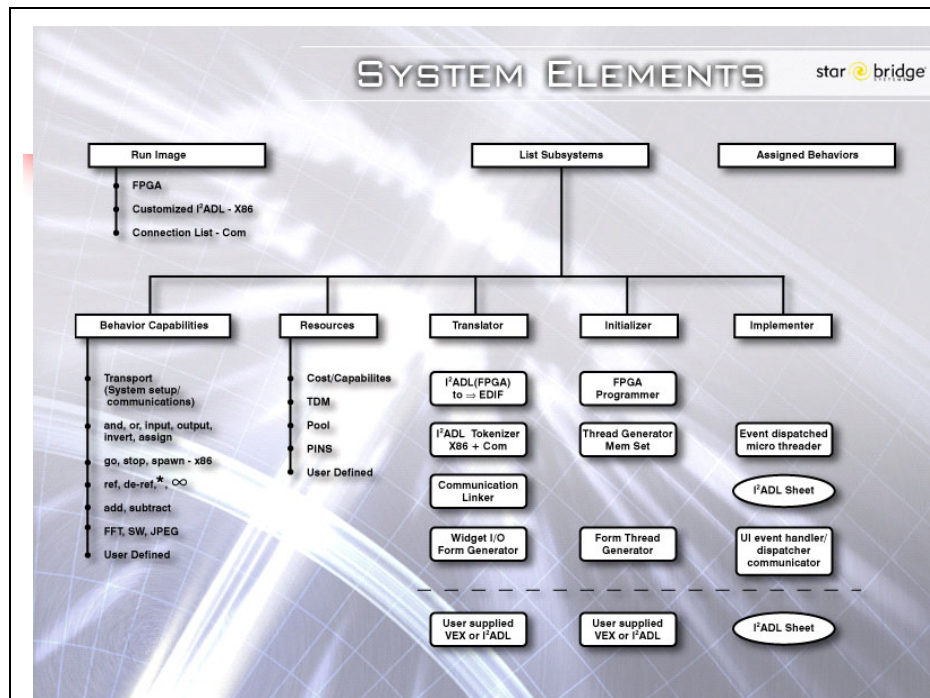


Figure 68



Libraries

- Application Builder Libraries MetaLIB
 - Math
 - Memory
 - I/O
 - Control
 - Logic Structure
 - Image Processing
 - Signal Processing
 - Data Compression

Figure 69



Libraries

- System Target Libraries
 - Emulation
 - Single Symmetric
 - 1 FPGA
 - Full Resource
 - Multiple FPGAs
 - Distributed over Multiple Boards or Systems

Figure 70



Libraries

- User Interface Component Libraries
 - Com/ActiveX
 - File I/O
 - Strings
 - Memory Management
 - Data I/O

Figure 71



Star Bridge Systems

- Bio-Informatics Application
 - Smith Waterm
- Air Force 1553 Interface Protocol
- Porting Viva to other FPGA systems
- Programming Hints

Figure 72



Smith Waterman Algorithm

- Search databases for sequences similar to a query sequence
- Dynamic programming to determine an optimal alignment
- Score is assigned for each character-to-character comparison
- Used to determine the position of matches

Figure 73



Smith-Waterman Algorithm

- The Smith-Waterman algorithm compares segments of all possible lengths (LOCAL alignments) and chooses whichever to maximise the similarity measure
- For every cell the algorithm calculates ALL possible paths leading to it. These paths can be of any length and can contain insertions and deletions

Figure 74

Smith-Waterman Algorithm

- Only works effectively when gap penalties are used

- In example shown

- match = +1
- mismatch = -1/3
- gap = -1+1/3k (k=extent of gap)

- Start with all cell values = 0

- Looks in subcolumn and subrow shown and in direct diagonal for a score that is the highest when you take alignment score or gap penalty into account

	C	A	G	C	C	U	C	G	C	U	U	A	G
A	0.0	1.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	1.0	0.0
A	0.0	1.0	0.7	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	1.0	0.7
U	0.0	0.0	0.8	0.3	0.0	0.0	0.0	0.0	0.0	1.0	1.0	0.0	0.7
G	0.0	0.0	1.0	0.3	0.0	0.0	0.7	1.0	0.0	0.0	0.7	0.7	1.0
C	1.0	0.0	0.0	2.0	1.3	0.3	1.0	0.3	2.0	0.7	0.3	0.3	0.3
C	1.0	0.7	0.0	1.0	3.0	1.7	?						
A													
U													
U													
G													
A													
C													
G													

$$H_{i,j} = \max\{H_{i-1,j-1} + s(a_i, b_j), \max\{H_{i-k,j} - W_k\}, \max\{H_{i,j-l} - W_l\}, 0\}$$

Figure 75

Smith-Waterman Algorithm

- Four possible ways of forming a path

For every residue in the query sequence

1. Align with next residue of db sequence ... score is previous score plus similarity score for the two residues
2. Deletion (i.e. match residue of query with a gap) ... score is previous score minus gap penalty dependent on size of gap
3. Insertion (i.e. match residue of db sequence with a gap) ... score is previous score minus gap penalty dependent on size of gap
4. Stop ... score is zero

Figure 76

Smith-Waterman Algorithm

Construct Alignment

- The score in each cell is the maximum possible score for an alignment of ANY LENGTH ending at those coordinates
- Trace pathway back from highest scoring cell
- This cell can be anywhere in the array
- Align highest scoring segment

GCC-UCG
GCCAUUG

	C	A	G	C	C	U	C	G	C	U	U	A	G
A	0.0	1.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	1.0	0.0
A	0.0	1.0	0.7	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	1.0	0.7
U	0.0	0.0	0.8	0.3	0.0	0.0	0.0	0.0	0.0	1.0	1.0	0.0	0.7
G	0.0	0.0	1.0	0.3	0.0	0.0	0.7	1.0	0.0	0.0	0.7	0.7	1.0
C	1.0	0.0	0.0	2.0	1.3	0.3	1.0	0.3	2.0	0.7	0.3	0.3	0.3
C	1.0	0.7	0.0	1.0	3.0	1.7	1.3	1.0	1.3	1.7	0.3	0.0	0.0
A	0.0	2.0	0.7	0.3	1.7	2.7	1.3	1.0	0.7	1.0	1.3	1.3	0.0
U	0.0	0.7	1.7	0.3	1.3	2.7	2.3	1.0	0.7	1.7	2.0	1.0	1.0
U	0.0	0.3	0.3	1.3	1.0	2.3	2.3	2.0	0.7	1.7	2.7	1.7	1.0
G	0.0	0.0	1.3	0.0	1.0	1.0	2.0	3.0	2.0	1.7	1.3	2.3	2.7
A	0.0	1.0	0.0	1.0	0.3	0.7	0.7	2.0	3.0	1.7	1.3	2.3	2.0
C	1.0	0.0	0.7	1.0	2.0	0.7	1.7	1.7	3.0	2.7	1.3	1.0	2.0
G	0.0	0.7	1.0	0.3	0.7	1.7	0.3	2.7	1.7	2.7	2.3	1.0	2.0
G	0.0	0.0	1.7	0.7	0.3	0.3	1.3	1.3	2.3	1.3	2.3	2.0	2.0

Figure 77

Smith-Waterman

Smith-Waterman

- Local alignments
- Residue alignment score may be positive or negative
- Requires a gap penalty to work effectively
- Score can increase, decrease or stay level between two cells of a pathway

Figure 78

Bio-Informatics Dilemma

- Sun Microsystems, Time Logic
 - Hardware Accelerators
 - 372,119 X 19,192 Comparisons
 - 41 hours 46 minutes
 - 1/10 of time without an accelerator
- Linux Cluster
 - 144 days of uninterrupted processing time
- Star Bridge Solution
 - Less than a second

Figure 79

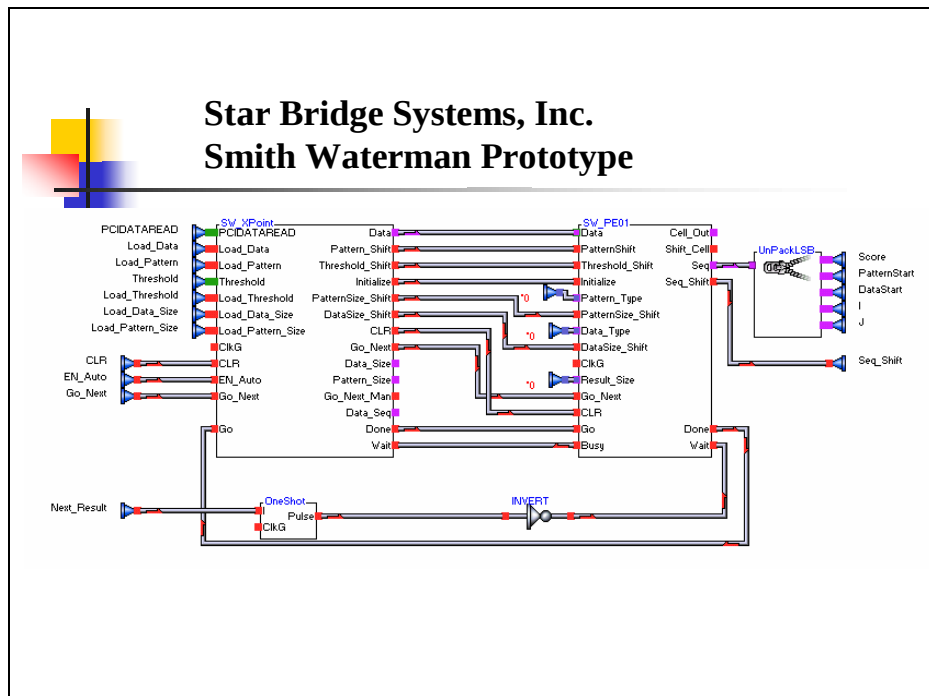


Figure 80

Main Smith Waterman Modules

- SW_Xpoint – FPGA which accesses PCI and User Interface busses. It gathers the data and presents it to SW_PE01, where calculations will take place.
- SW_PE01- FPGA where Smith Waterman iterations are actually implemented.

Figure 81

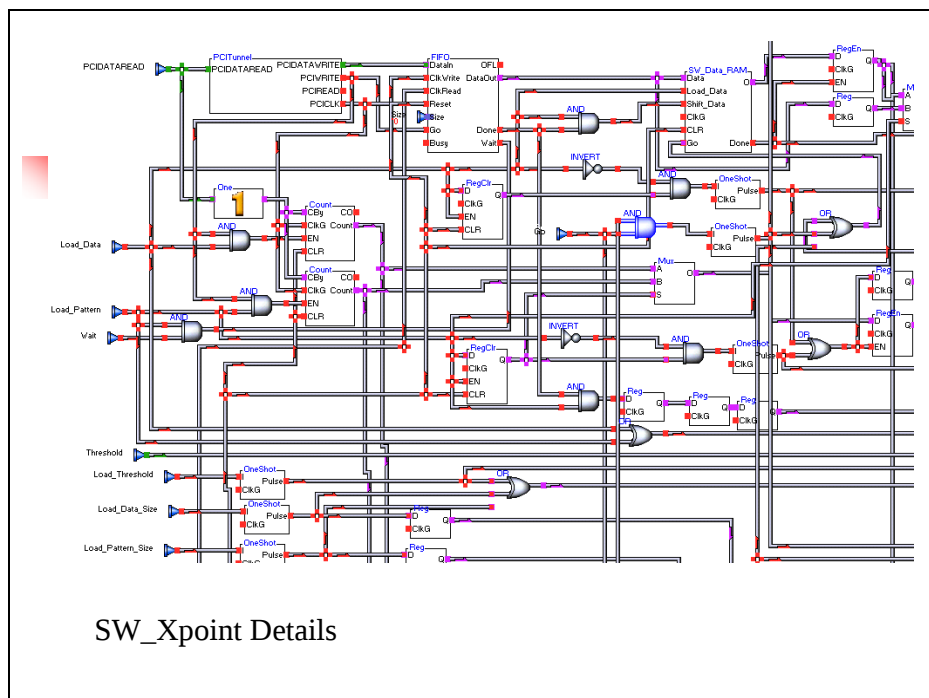


Figure 82

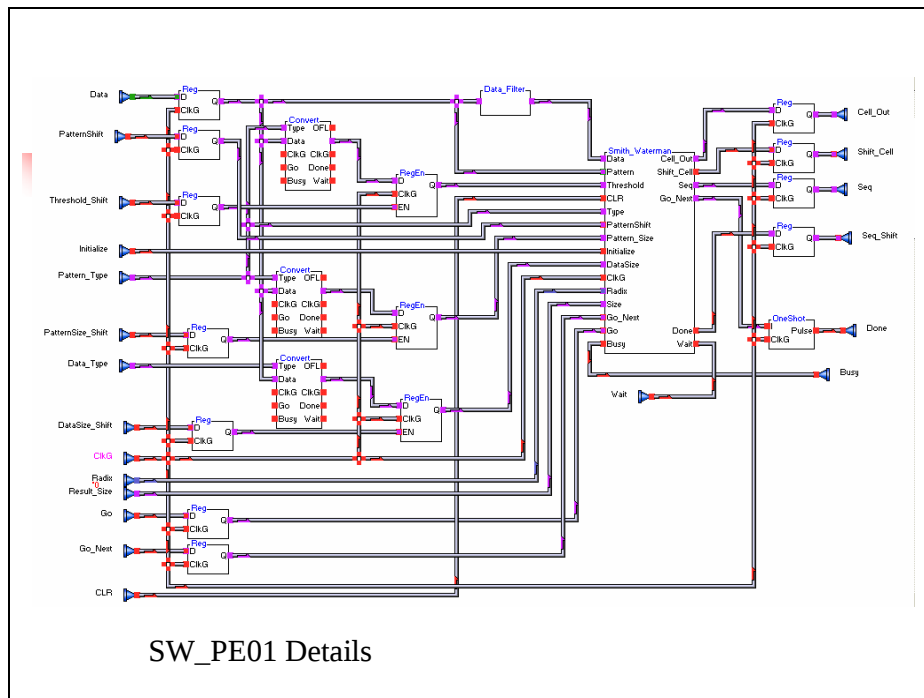


Figure 83

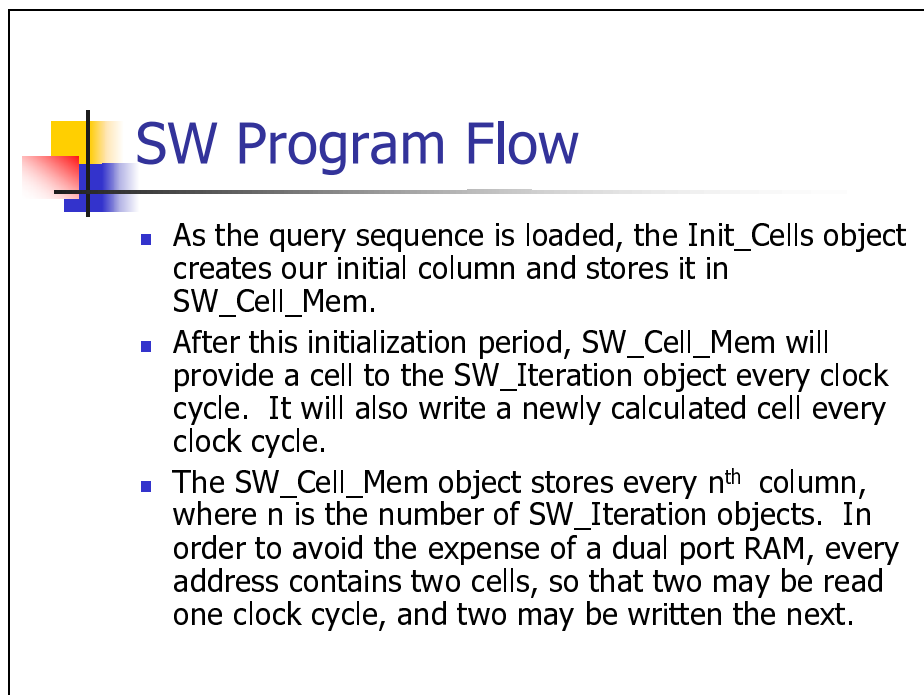


Figure 84



Smith Waterman Cells

- There are as many cells as there are characters in the query sequence.
- The array of cells represent a column of the scoring matrix.
- The initial (zero) column is initialized and stored into the cell memory object, SW_Cell_Mem.

Figure 85



Cell Contents

- Each cell contains the following four parameters:
- Pattern – a character from the query sequence
- Score – the score of this cell in the current i,j position
- PatternStart – the position in the query sequence from which the score was calculated
- DataStart – the position in the reference sequence from which the score was calculated

Figure 86



		0		2	3	4	5	6	7	8	9
		0	G	A	A	G	A	A	G	C	G
score		0	0	0	0	0	0	0	0	0	0
patst	0	0	0	0	0	0	0	0	0	0	0
datast	0	0	1	2	3	4	5	6	7	8	9
score		0	0	1	1	0	1	1	0	0	0
patst	A	1	1	0	0	1	0	0	1	1	1
datast	1	0	1	1	2	4	4	5	7	8	9
score		0	0	1	2	0	1	2	0	0	0
patst	A	2	2	1	0	2	1	0	2	2	2
datast	2	0	1	1	1	4	4	4	7	8	9
score		0	1	0	0	3	0	0	3	0	1
patst	G	3	2	3	3	0	3	3	0	3	2
datast	3	0	0	2	3	1	5	6	4	8	8
score		0	0	0	0	0	2	0	0	4	0
patst	C	4	4	4	4	4	0	4	4	0	0
datast	4	0	1	2	3	4	1	6	7	4	4
score		0	1	0	0	1	0	1	1	0	5
patst	G	5	4	5	5	4	5	0	4	5	0
datast	5	0	0	2	3	3	5	1	6	8	4

Figure 87

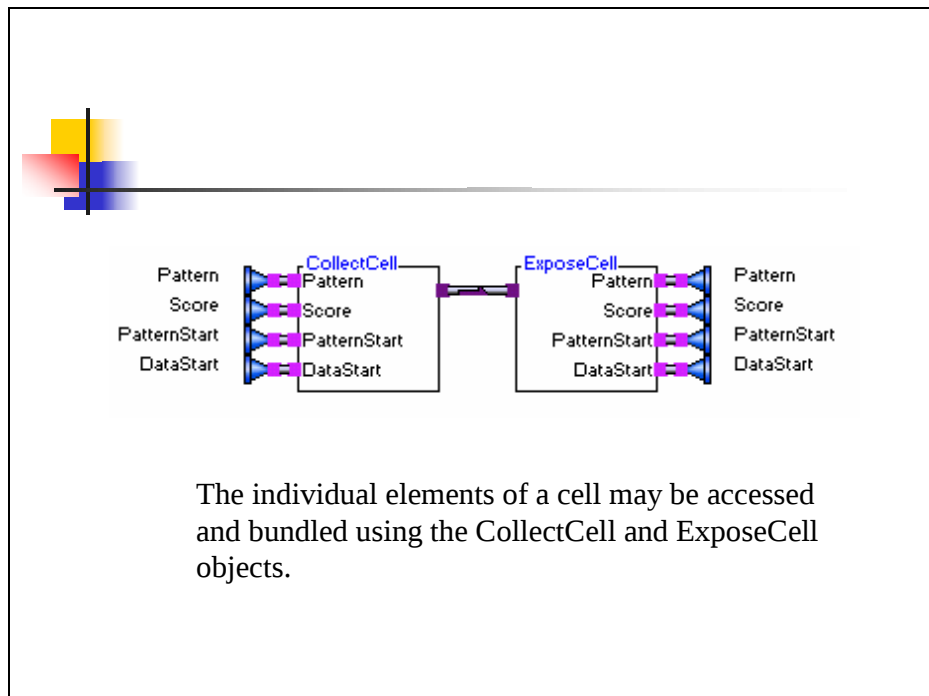


Figure 88



Cell Data Types

- The size of all data elements may be adjusted depending on usage:
- Pattern – contains as many bits as needed to encode characters from the sequences – 4 bits for genes.
- Score and PatternStart – must be the same size and be large enough to encode the number of entries in the query sequence
- DataStart – will be the largest data set as it must be able to encode any position in the reference sequence.

Figure 89



Right Size for the Job

- Because all the parameters are adjustable in their size, less circuitry is needed to calculate matches in smaller sequences.
- Smaller sequences may for this reason utilize more parallelism.

Figure 90

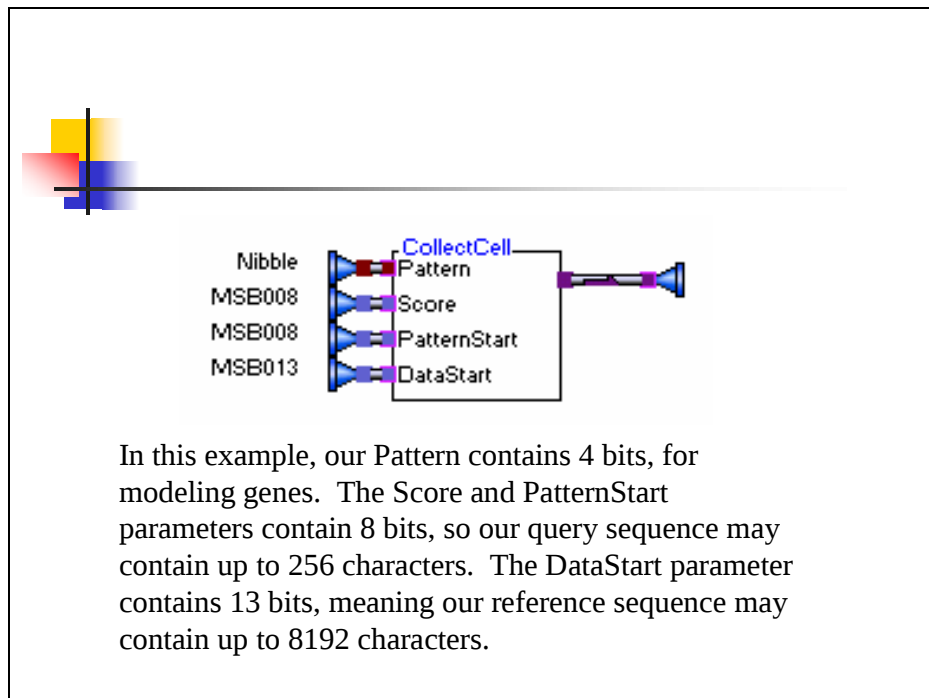


Figure 91

Parallelism

- If a given hardware system has enough physical resources to accommodate n SW_Iteration objects, the Smith Waterman program may operate on n columns in parallel.
- Hence n cells are computed every clock cycle.
- Conservative estimates place 150 SW_Iteration objects in each Virtex II 6000

Figure 92



HC-62

- An HC 62 has the bandwidth to pass cells between 8 FPGAs, allowing for 720 parallel SW_Iteration objects.
- At a conservative 100Mhz system clock speed, this gives $100,000,000 * 720 =$
- 72 billion Smith Waterman steps/second

Figure 93



Smith-Waterman Solutions

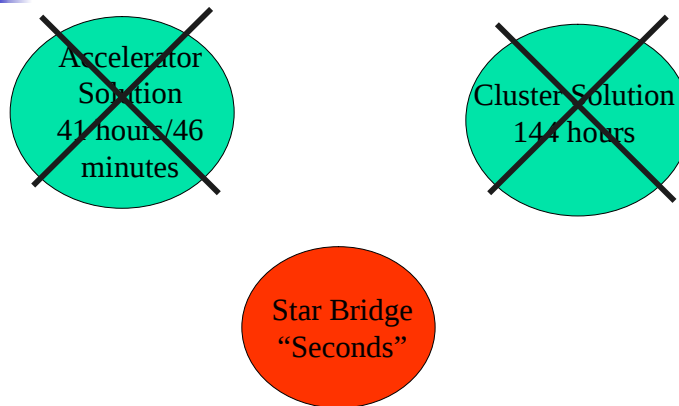


Figure 94



1553 Bus Controller

- IP Library element
- Reduce costs
- Reduce number of parts
- Eliminate parts obsolescence problems
- Increase flexibility
 - Multiple channels on a chip
 - Programmability for usage
 - Bus Controller, Remote Terminal, or Bus Monitor

Figure 95



1553 Protocol

- Polymorphic
 - Overloading
 - Recursion
 - Synthesis time resolution
- Polymorphic 1553
 - Command word
 - Data word
 - Status word
 - Programmability in encoding, length of data field, construction of the sync waveform, or the speed the data is transmitted

Figure 96



1553 Protocol

- Viva Implementation
 - 150 CLBs Vs. 2800 logic elements from a competitor
 - Short development time

Figure 97



Target Viva to XSV-800

- System Description
 - FPGA System
 - Clock System
 - Parallel I/O System
 - Parallel Input Behavioral Communications
 - Viva Port I/O Object
- Programming the FPGA
- Viva Programming

Figure 98

Target Viva to XSV-800

- Programming the FPGA
 - Implementation System
 - Overloads board initializer program
 - Viva spawns file progxsv.vex
 - Calls a DOS application provided by board manufacturer
 - Spawns Xsload, also provided by board manufacturer
 - Viva Programming
 - Program development
 - EDIF file generation
 - Xilinx place and route
 - Spawns the progxsv.vex program

Figure 99

XSV-800 Board

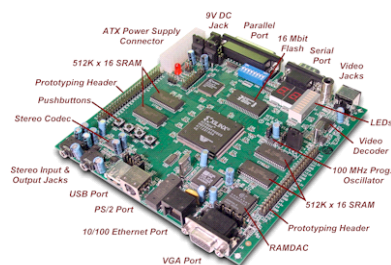
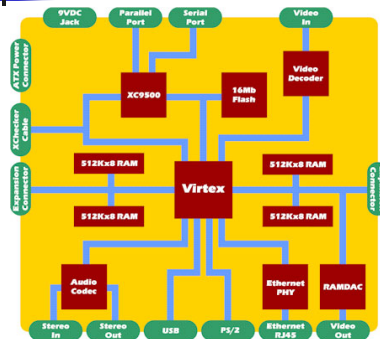


Figure 100



Star Bridge Systems Hypercomputing

- Greatly Reduces Development Cycle
- Reduces Time to Market
- Parallel and Reconfigurable Computing
- Greater Programming Flexibility
- Lower Cost
- Lower Power Consumption
- Smaller Size
- Breakthrough in Performance and Speed

Figure 101



Recurive Examples

- Learn by taking apart examples and library elements.
- ADC example
- Have Fun Programming in Viva

Figure 102

DARPA'S NEW COGNITIVE SYSTEMS VISION

Zachary Lemnios
DARPA/Information Processing Technology Office
Arlington, VA

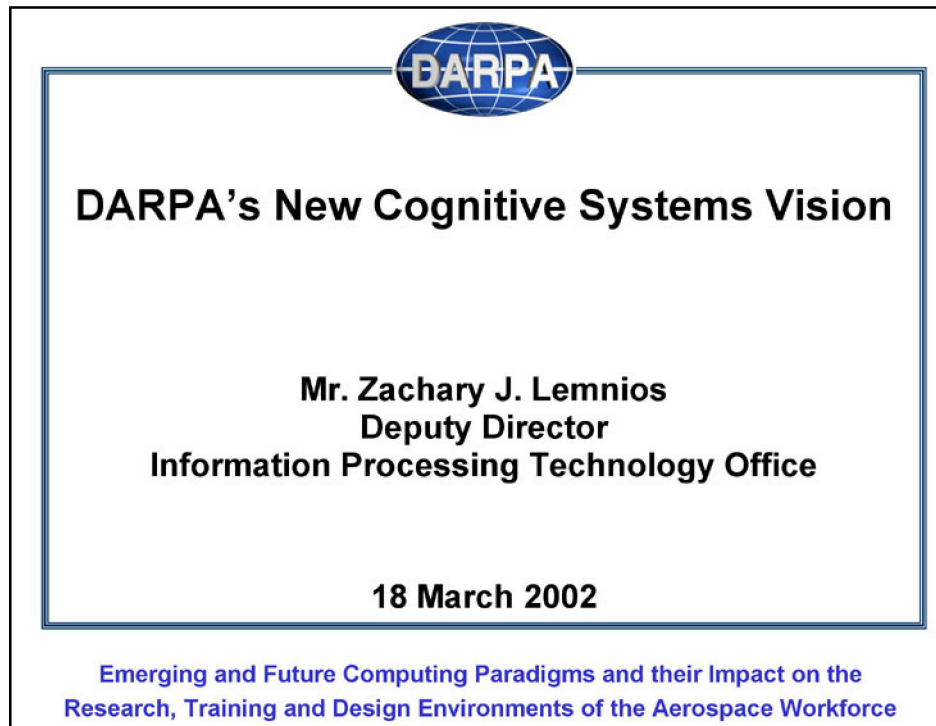


Figure 1

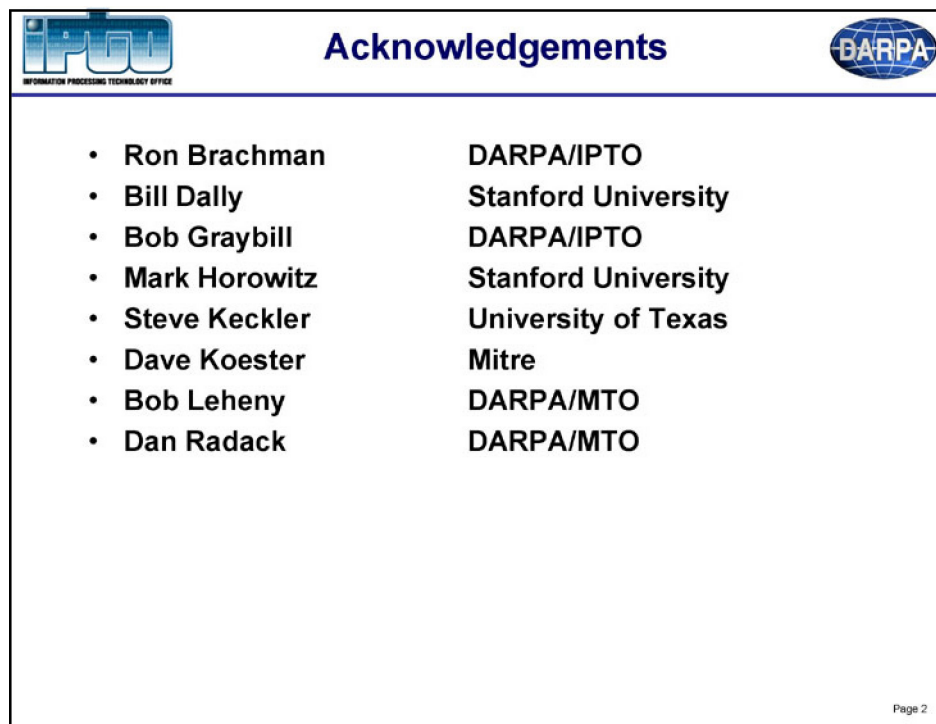


Figure 2



INFORMATION PROCESSING TECHNOLOGY OFFICE

Agenda



➔ • **Introduction**

- DoD System Challenges
- Motivation for Cognitive Computing

• **Technology Trends**

- Device Performance
- High Performance Computing

• **Cognitive Systems Vision**

• **Summary**

Page 3

Figure 3

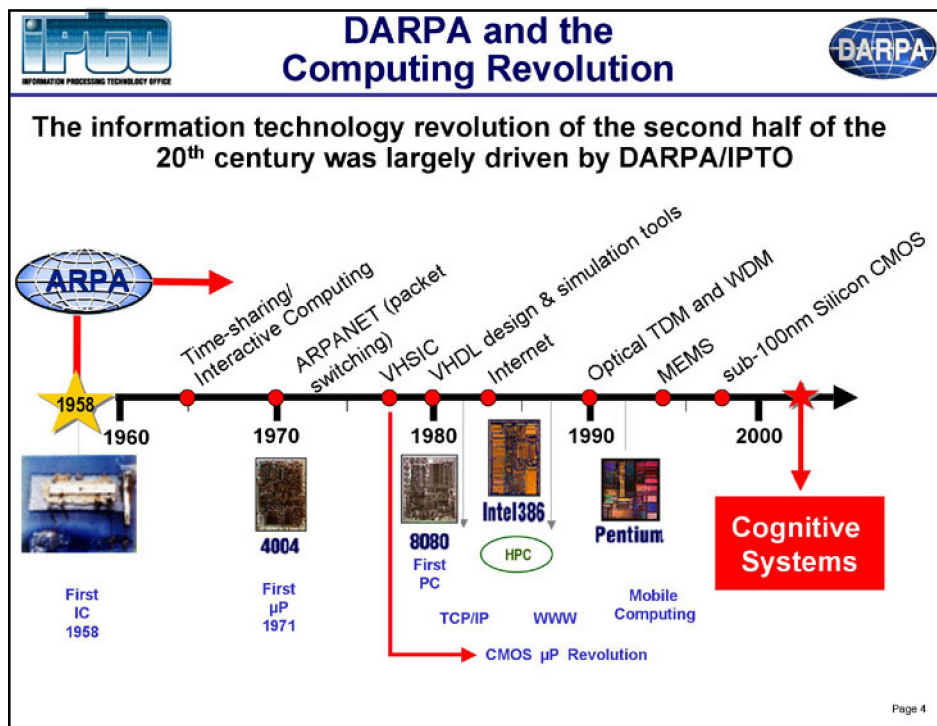


Figure 4

A NEW CLASS OF AUTONOMOUS SYSTEMS

This chart illustrates the challenge for autonomous systems, to provide increasing performance to enable new capabilities in environments of increasing complexity.

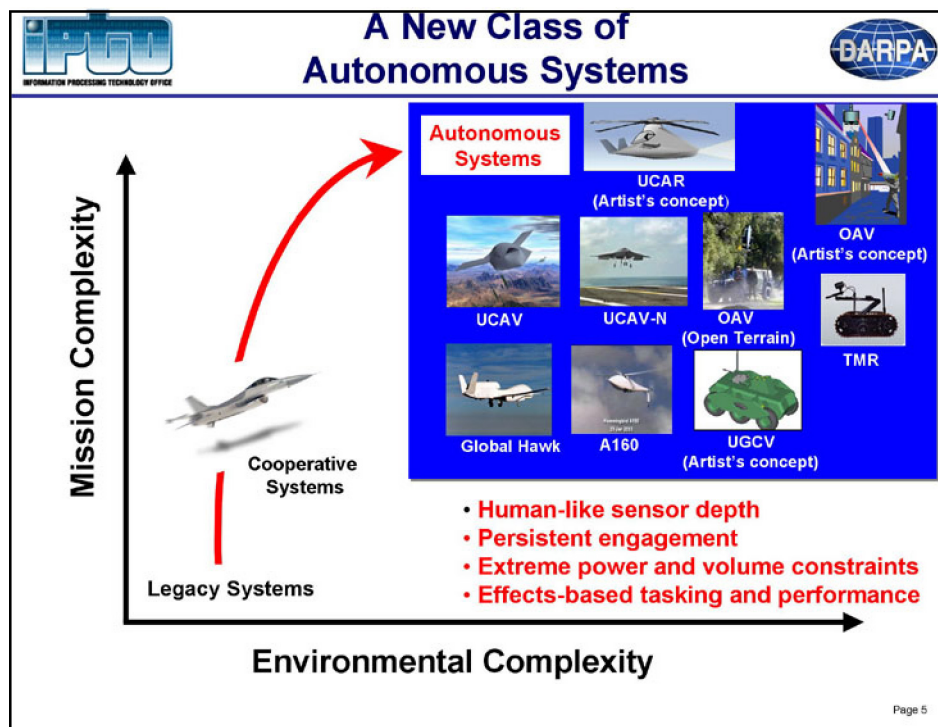


Figure 5

THE CHALLENGE OF COMPLEXITY

The major focus for the newly formed IPTO is the development of cognitive systems. While computational performance has been increasing due to Moore's Law, the productivity and effectiveness of these systems are not increasing at the same rates. Development of cognitive systems will enable systems to become more usable; more flexible in application; less vulnerable to attacks and more robust in detecting and recovering; all while remaining cost-effective. Cognitive systems are enabled by a firm foundation in the underlying science and mathematics in algorithms and information assurance, which is then embodied in robust software and executed on advanced hardware. Future IPTO programs will also focus on the critical aspects of autonomous perception; knowledge representation and reasoning; machine learning; and communications and interactions.

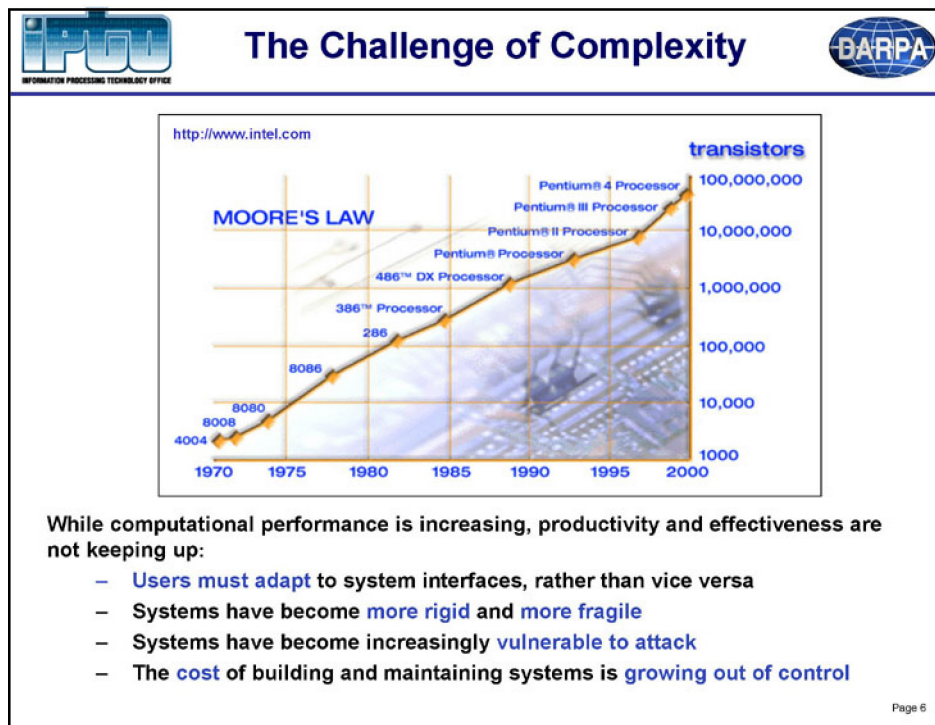


Figure 6

DARPA/IPTO COGNITIVE SYSTEMS

The major focus for the newly formed IPTO is the development of cognitive systems. While computational performance has been increasing due to Moore's Law, the productivity and effectiveness of these systems are not increasing at the same rates. Development of cognitive systems will enable systems to become more usable; more flexible in application; less vulnerable to attacks and more robust in detecting and recovering; all while remaining cost-effective. Cognitive systems are enabled by a firm foundation in the underlying science and mathematics in algorithms and information assurance, which is then embodied in robust software and executed on advanced hardware. Future IPTO programs will also focus on the critical aspects of autonomous perception; knowledge representation and reasoning; machine learning; and communications and interactions.

**DARPA/IPTO**
Cognitive Systems

- **DARPA IPTO will create a new generation of cognitive computational and information systems with capability to:**
 - **reason**, using substantial amounts of appropriately represented knowledge
 - **learn** from their experience so that they perform better over time
 - **explain** themselves and be told what to do
 - **be aware** of their own capabilities and reflect on their own behavior
 - **respond** robustly to surprise

Systems that know what they're doing

Page 7

Figure 7

WHY NOW?

Cognitive technology (from AI) is working in bits and pieces, ranging from large-scale knowledge bases to machine learning in support of data mining

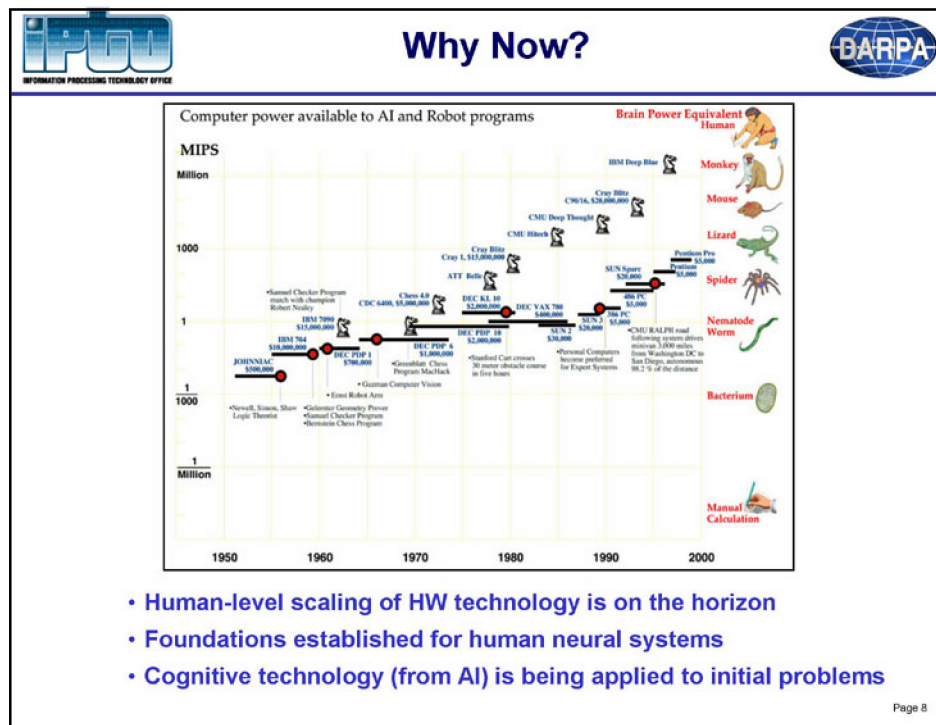


Figure 8



INFORMATION PROCESSING TECHNOLOGY OFFICE

Agenda



- Introduction
 - DoD System Challenges
 - Motivation for Cognitive Computing
- ➔ • Technology Trends
 - Device Performance
 - High Performance Computing
- Cognitive Systems Vision
- Summary

Page 9

Figure 9

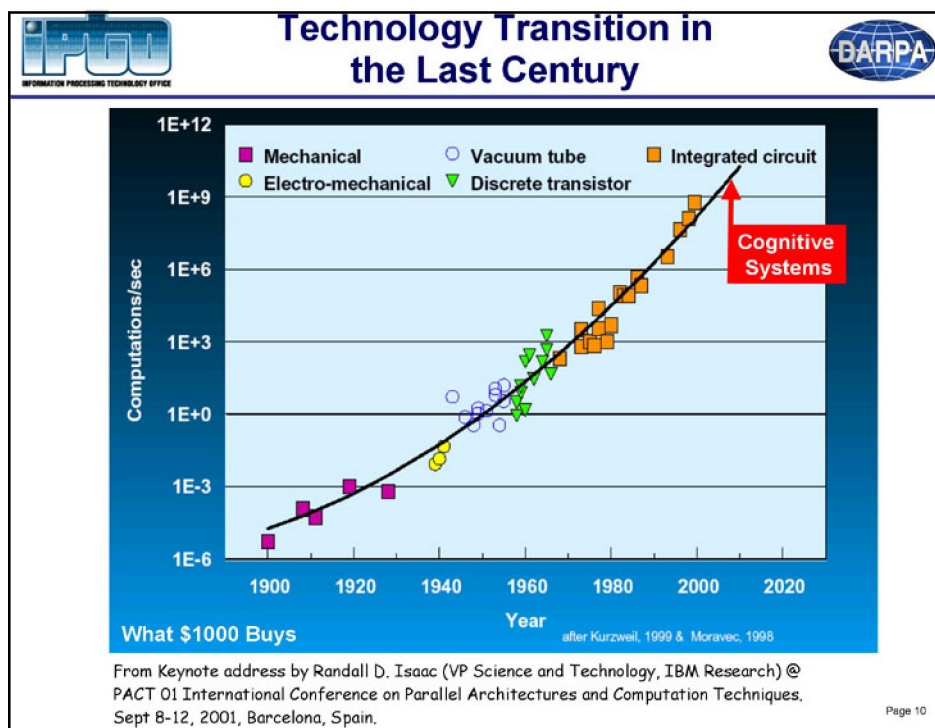


Figure 10

MEMORY WALL IN GROWING

Based on the SIA roadmap projections, as CMOS feature sizes decrease along with a clock frequency increase, the actual memory access times in clock cycles actual increase. Memory access time are also a factor of the memory size or capacity due to increased wiring delays. The memory wall only increases as memory is moved further and further away from the processor core. (Off chip or off the board) The trend to date has been to move memory on board the chip but there is only so much that can be accomplished as indicated by this graph. One solution is to develop Processor-in-Memory (PIMS) with many small memory tiles local to small processor cores.

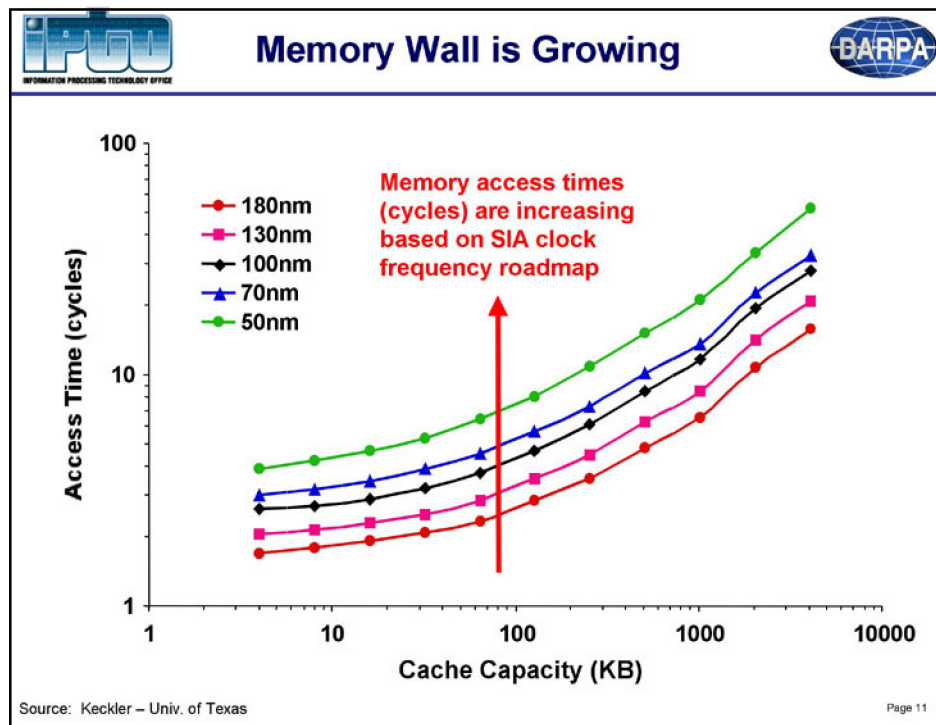


Figure 11

SIA ROADMAP IMPACT ON COMPUTER ARCHITECTURES

This is a vivid example of why computer architectures will not scale with CMOS feature size. The spatial extent of a clock (logic functions) has shrunk from an entire die to a very small region within a die. New innovative architectures, packaging and interconnect are required to efficiently take advantage of the increased number of transistors per MM. Spatial locality is very critical. Note: Wire delays are not decreasing at the same rate as the transistor delays.

SIA Roadmap – Another interesting trend from the physics of signal propagation is illustrated here. When clocks were slow, electrical signals could easily traverse across the chip in the time to settle between subsequent timing edges. As clock speeds have increased, the physical distance that a signal can travel and settle between subsequent timing edges has decreased. This trend will continue, even with the introduction of repeaters. It is expected that changes in circuit implementation, such as a greater exploitation of asynchronous communication will begin, and that these changes will also affect the performance of various new architectures.

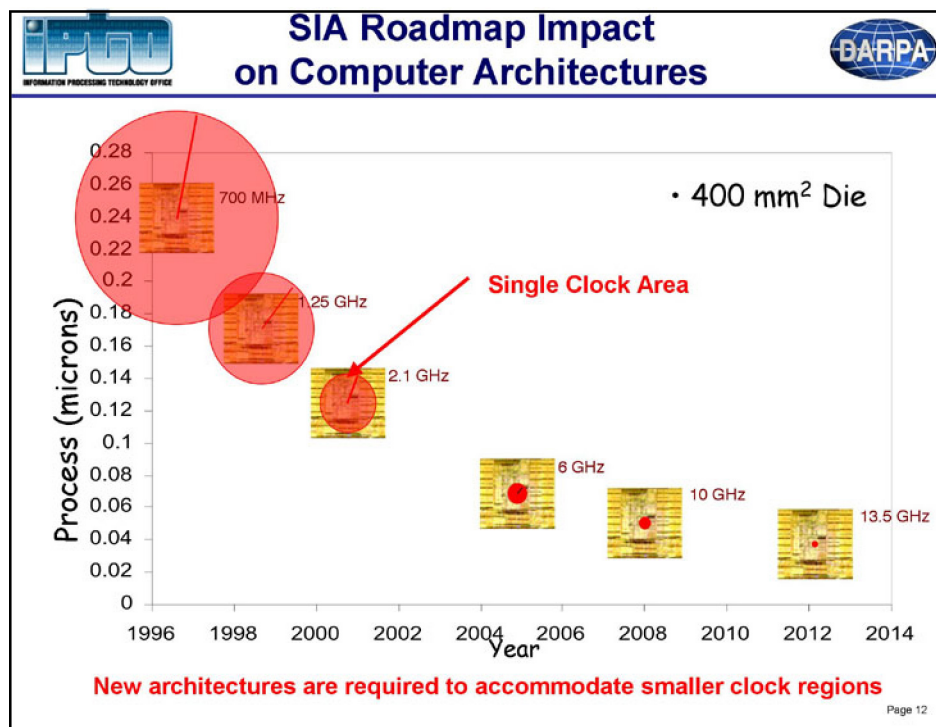


Figure 12

NOVEL ARCHIECTURES ARE NEEDED TO EXTEND PRODUCTIVITY

To date the programming model has hidden the architectural (tricks) techniques used to achieve Moores Law.

Conventional processor scaling or peak performance improvements are going to slow down due to limited improvement in the number of instructions per cycle and end of pipelining advantages (number of implemental gates between clock cycles).

The realization of these techniques to-date have resulted in enormous increased complexity, chip design cost/schedule, and chip software costs. To the point where only a few large elite companies can now afford to develop new families of computers.

The challenge is to now exploit and develop new solutions to make systems more productive!

Over the past 20 years (and in fact since the first microprocessor in 1971), microprocessor performance has been improving at a rate of approximately 52% per year. This exponential increase in performance has come from three sources: (1) improvements in clock rates from innovations in silicon processing technology that have made transistors smaller and faster – 19%/year, (2) improvements in clock rates resulting from deepening pipelines and reducing the number of levels of logic per clock – 9%, and (3) improvements in microarchitectures including multi-instruction issue and out-of-order execution – 18%. Modern designs have nearly exhausted the benefits of pipelining, and conventional architectures are struggling to sustain even one instruction per cycle. Without further innovations, performance improvements will at best only match the rate of improvement due to further process technology innovations, which is projected to continue at 19% per year. While microprocessors have sustained performance improvements of 52%/year, fabrication technology has actually provided a much higher growth rate in potential capability. When accounting for increased transistor counts and faster transistor switching speeds, the capability of microprocessor-scale integrated circuits has been improving at 74%/year. Until now, the differential between the 74% and 52% rates has resulted in only a factor of 30 of untapped performance potential. However, with only 19% per year projected in the future, the differential is expected to increase to a factor of 30,000 by 2020. This quantity represents a tremendous opportunity for novel architectures to help bridge the performance gap and to enable future computer systems to solve increasingly complex and important problems.

EMBEDDED COMPUTING PERFORMANCE REGIONS

FOM2.PPT, James C. Anderson, 6/5/00

Upper diagonal line: 6U VME limit is $6A@5V$ & $1A@+/-12V = 54W$ per slot (6.3x9.2x0.8" slot-to-slot spacing) @60C, "System Packaging Products," Carlo Gavazzi, Inc., Mupac Business Unit. SEM-E (MIL-STD-1389 & IEEE-Std-1101.4-1993) is 12W conservative, 24W typ & 48W peak (5.88x6.41x0.6" slot-to-slot spacing).

Projected requirements for processor subassemblies of selected DARPA projects: Space-Based Radar ca. 2008 (1100 GOPS, 20W, 0.05 cu ft), Uninhabited Aerial Vehicle ca. 2005 (710 GOPS, 400W, 2 cu ft), Soldier's Radio for Small Unit Operations ca. 2002 (13 GFLOPS, 0.5W, 0.022 cu ft)

Graybill Notes: This vg highlights the relative upper boundaries for the three original and the new PCA class of processing options in terms of Computational efficiency (GOPS/Watt) and Computational Density (MOPS/cmsq). There is a key third dimension that is not shown that would highlight the great variation in efficiency as a function of kernel types

With the advent of Polymorphous Computing technology a new class of computing options will now be available for embedded computing. Polymorphous Computing offers almost the density of VLSI but with the programmability of conventional computers. In addition the architecture or virtual machine realization may be changed dynamically in response to mission processing/threat requirements. Agile virtual selection of computer types may now be done during the mission instead of locking into a pre-determined mix of processor types during platform development.

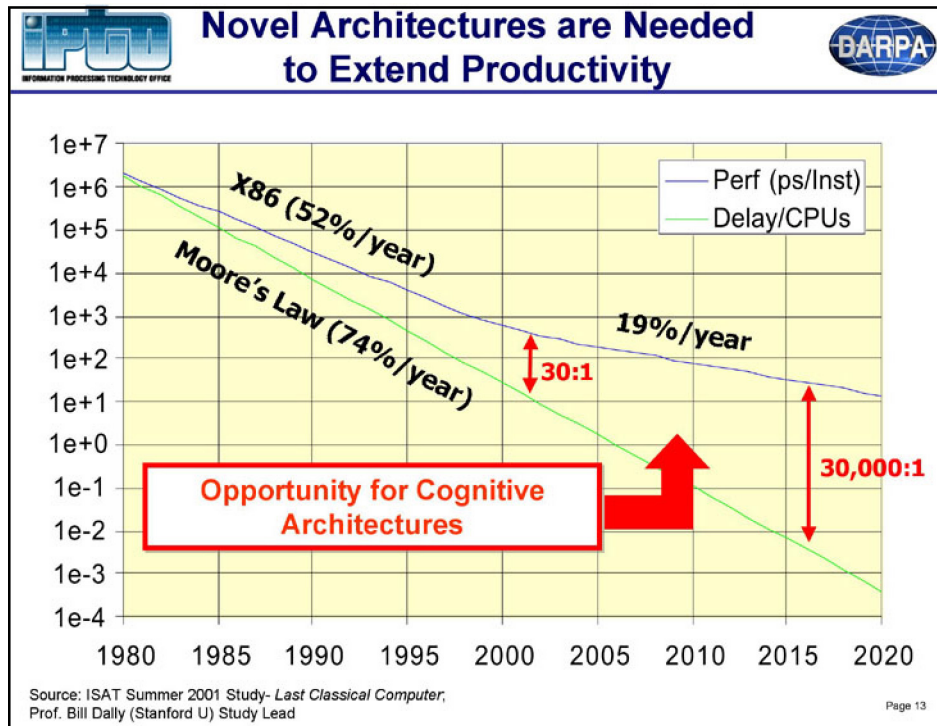


Figure 13

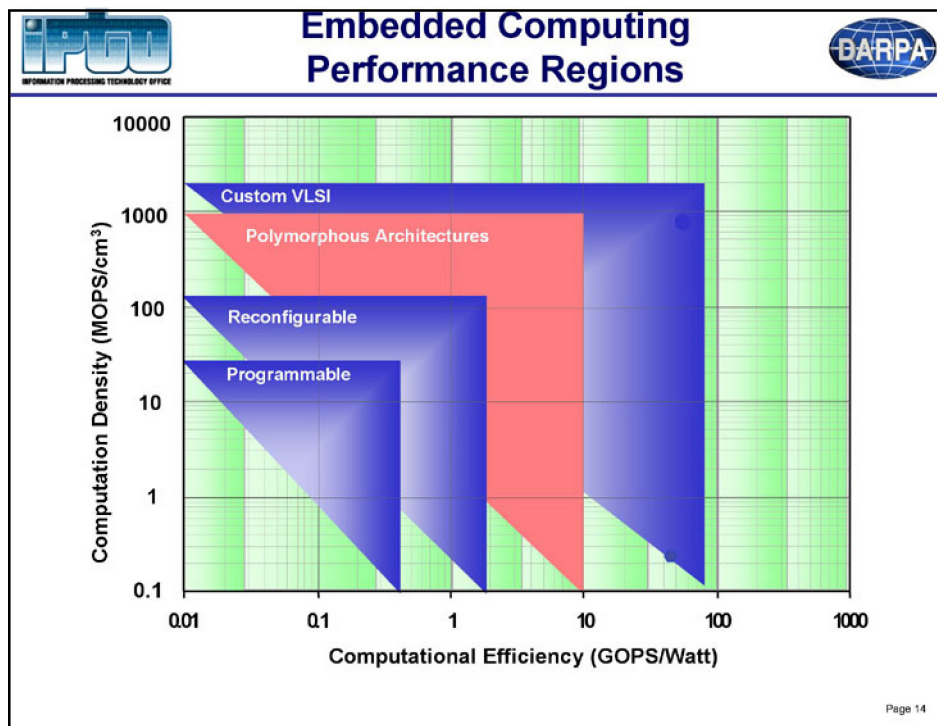


Figure 14

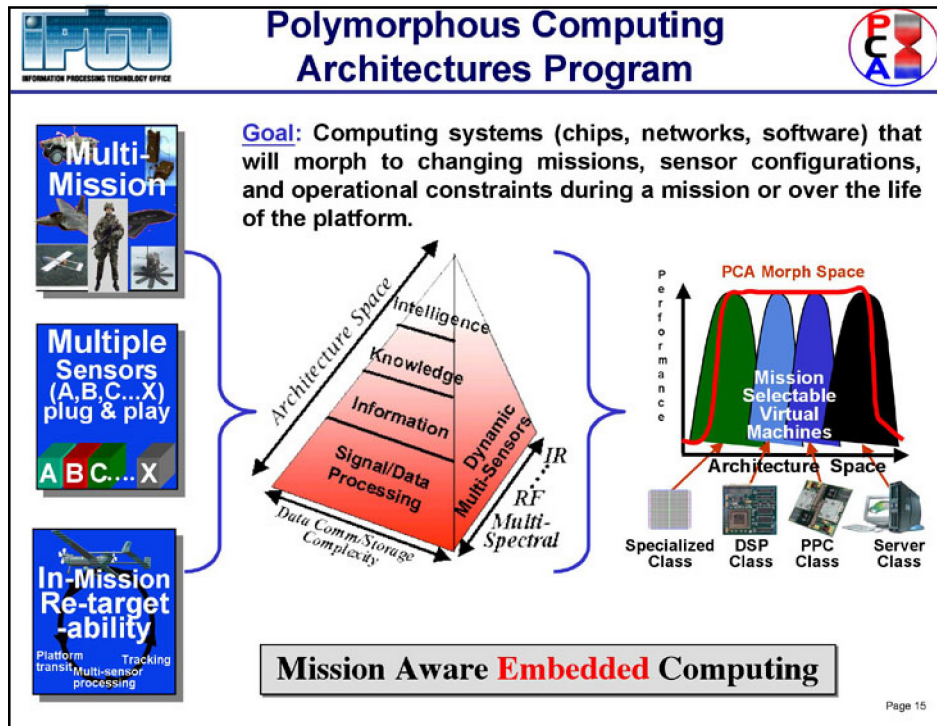


Figure 15

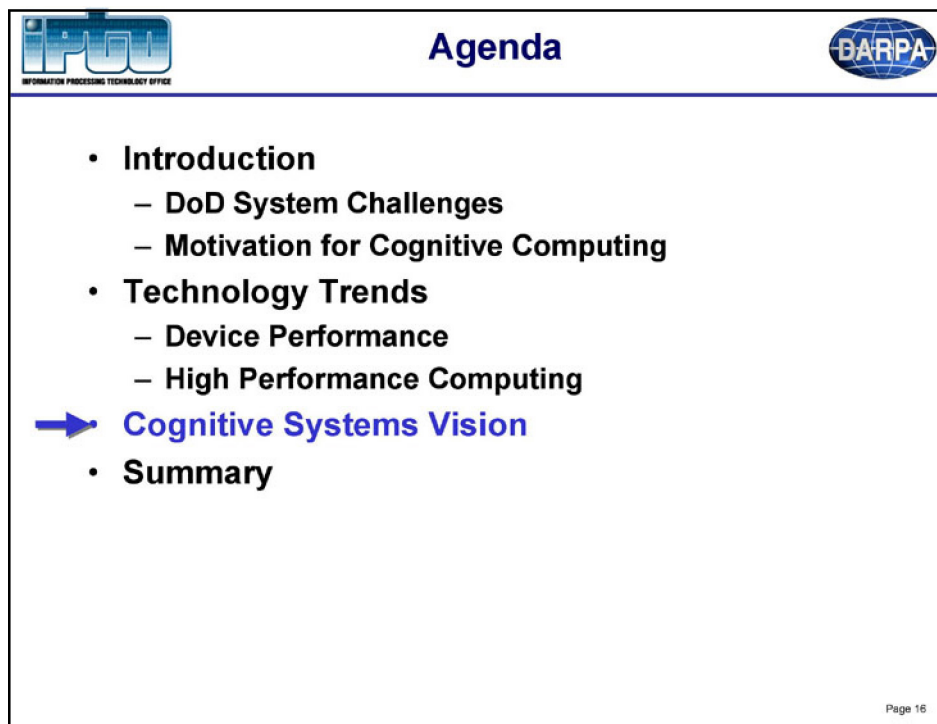


Figure 16

DARPA/IPTO COGNITIVE SYSTEMS

The major focus for the newly formed IPTO is the development of cognitive systems. While computational performance has been increasing due to Moore's Law, the productivity and effectiveness of these systems are not increasing at the same rates. Development of cognitive systems will enable systems to become more usable; more flexible in application; less vulnerable to attacks and more robust in detecting and recovering; all while remaining cost-effective. Cognitive systems are enabled by a firm foundation in the underlying science and mathematics in algorithms and information assurance, which is then embodied in robust software and executed on advanced hardware. Future IPTO programs will also focus on the critical aspects of autonomous perception; knowledge representation and reasoning; machine learning; and communications and interactions.

**DARPA/IPTO
Cognitive Systems**

- **DARPA IPTO will create a new generation of cognitive computational and information systems with capability to:**
 - **reason**, using substantial amounts of appropriately represented knowledge
 - **learn** from their experience so that they perform better over time
 - **explain** themselves and be told what to do
 - **be aware** of their own capabilities and reflect on their own behavior
 - **respond** robustly to surprise

Systems that know what they're doing

Page 17

Figure 17

ANATOMY OF A COGNITIVE AGENTS

This chart illustrates one possible architecture for realizing a cognitive computational system. This diagram shows the relationships and connectivity among the 3 major processes that are usually associated with cognition. In addition, the relationship between these processes and the machine's perception (sensors), action (effectors), and the environment.

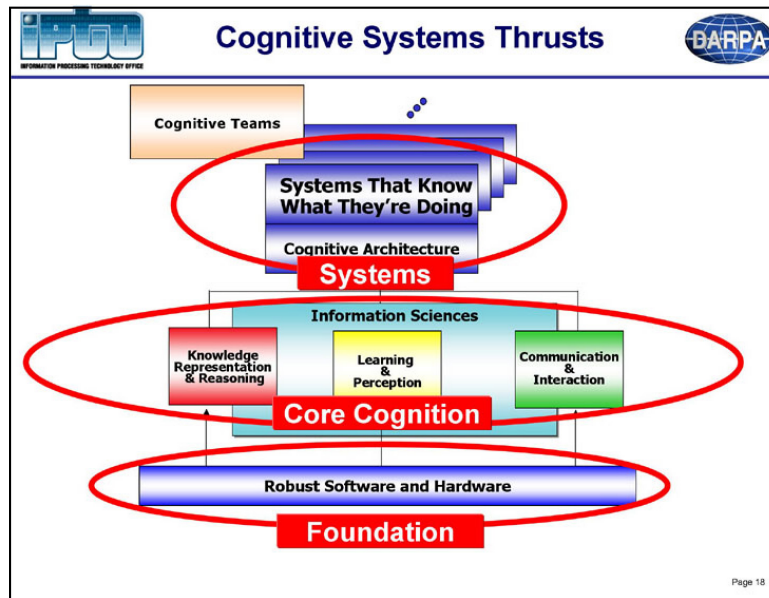


Figure 18

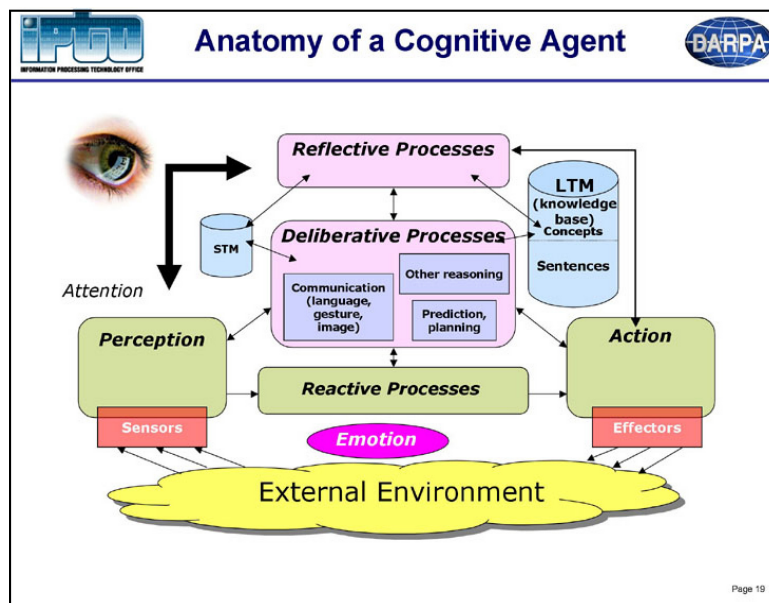


Figure 19



INFORMATION PROCESSING TECHNOLOGY OFFICE

Initial Challenge Context



Persistent, personal partner/associate systems

- Learn from experience
- Learn what you like and how you operate
 - 4 by observation
 - 4 by direct instruction or guidance, in a natural way
- Imagine possible futures, anticipate problems and needs
- Omnipresent / always available

Examples

- Commander's (C2) assistant
- (Intelligence) Analyst's associate
- Personal executive assistant/secretary
- Disaster response captain's "RAP" (robot/agent/person) team

Page 20

Figure 20



INFORMATION PROCESSING TECHNOLOGY OFFICE

Focal Challenge Context



An Enduring, Personalized, Cognitive Assistant





Radar O'Reilly:

Anticipated
Planned
Reasoned
Advised
Acted

and, never failed...

Page 21

Figure 21



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





RADAR O'REILLY

- Uses *knowledge* of the domain, task
- Exhibits *purposeful perception*
- Can *imagine* possible futures
- Can *decide* what to do (prioritize) and act
- *Learns*, including by observing partner
- Can be *advised* and *guided*, and can *explain*
- Must know how to *cooperate*
- *Multi-modal*, broad-spectrum interaction
- Available everywhere - *omnipresent*
- Must be *trustworthy*
- Must *learn continuously*
- Able to *survive*, operate through problems

ENDURING, PERSONALIZED, COGNITIVE ASSISTANT

Page 22

Figure 22

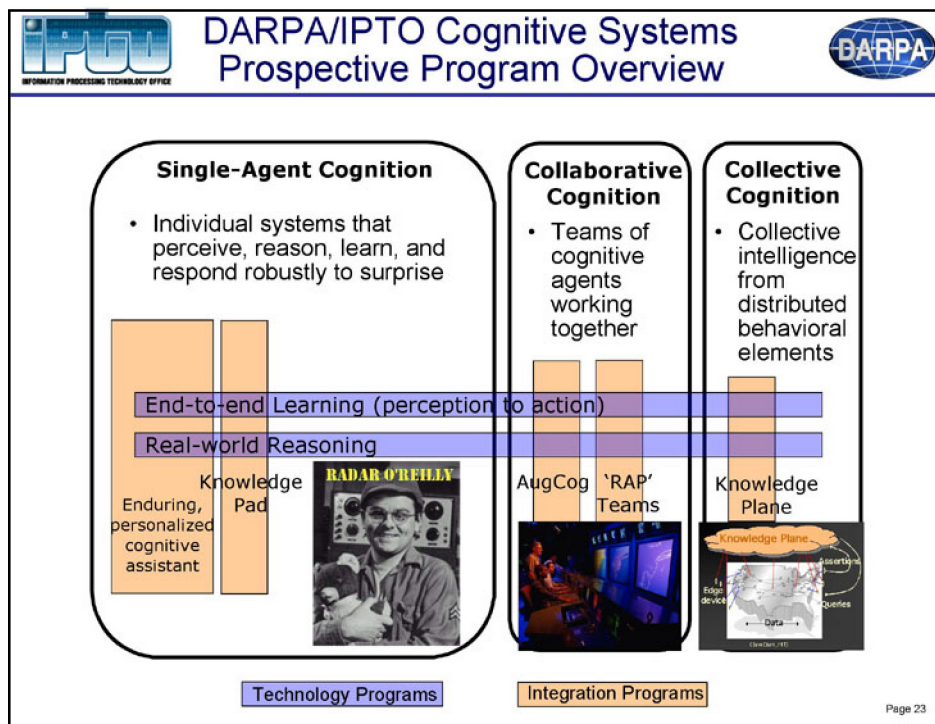


Figure 23



Summary

Cognitive Systems will be the Next Revolution in Computing

Cognitive Systems know what they are doing:

- They can reason
- They can learn from their experience
- They can explain themselves
- They are aware of their own capabilities
- They can respond robustly to surprise

Send us your best ideas:

http://www.darpa.mil/ipto/Solicitations/PIP_02-21.html

•Take a tour as a DARPA Program Manager

rbrachman@darpa.mil	(703) 696-2264
zlemnios@darpa.mil	(703) 696-2234

Page 24

Figure 24



Backup

Page 25

Figure 25

ASYMMETRIC ADVANTAGE ENABLED BY INFORMATION SUPERIORITY

“Not long ago, a prime contractor was one that built a jet or ship. Now, these vehicles are simply “platforms” for sensors and information systems.”

*Washington Post; Analysis of NG bid for
TRW, Feb. 23, 2002*

Asymmetric Advantage – Information superiority gives US forces an asymmetric advantage over our adversaries. This chart shows a possible future battlespace environment with a representative threat and a number of individual high performance platforms and capabilities that can be brought to bear against all threats. In addition, these platforms will also be networked together to enable instant sharing of information. The network will close the sensor-to-shooter latency for the classes of dynamic targets that are difficult to detect by conventional means. Advances in electronics and in computation will enable us to see farther with greater clarity than our enemies could ever imagine.

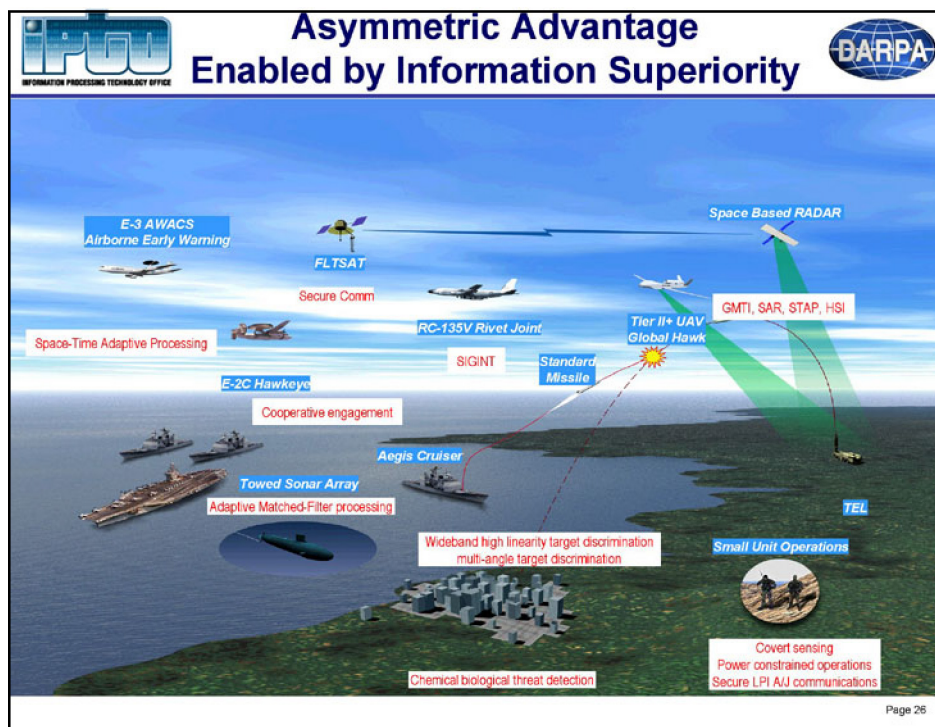


Figure 26

INFORMATION NETWORKS HAVE REVOLUTIONIZED PLATFORMS

Information Networks Will Revolutionize Platforms – Networks will also be used at the platform scale to revolutionize their capabilities. On the upside, these now networked systems will be more survivable, more lethal, more adaptable, and more agile than before. On the downside though is complexity. Not only are there issues relating to physical interconnections and signal routing, but there is also a massive growth in the number of lines of software code necessary for managing and exploiting the information on the network. Integration costs are also escalating. Actual hardware costs can be a fraction of the total, while software and integration/test are accounting for a major part of the “fly-away” cost of a platform.

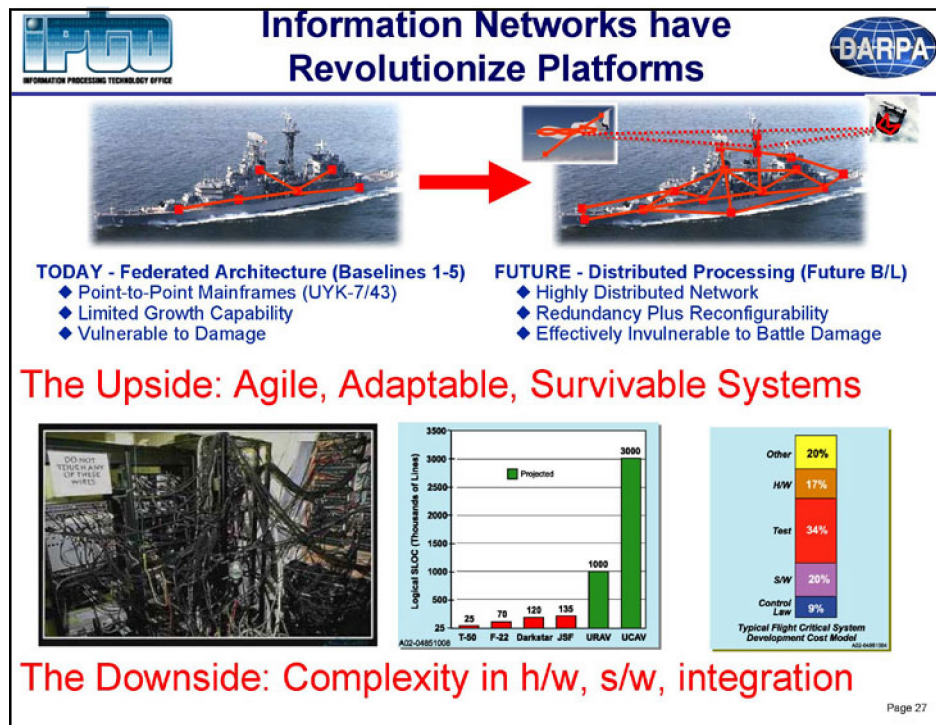


Figure 27

THE RESULT WILL ENABLE A REVOLUTION IN CAPABILITY FOR DOD

The Result Will Enable a Revolution in Capability for DoD – Potential threats are increasing and human analysts are overwhelmed by sensor data. Adaptive and intelligent data-fused sensors will enhance system performance against a changing threat environment. Cognitive information exploitation will provide the knowledge and information from fused sensors.

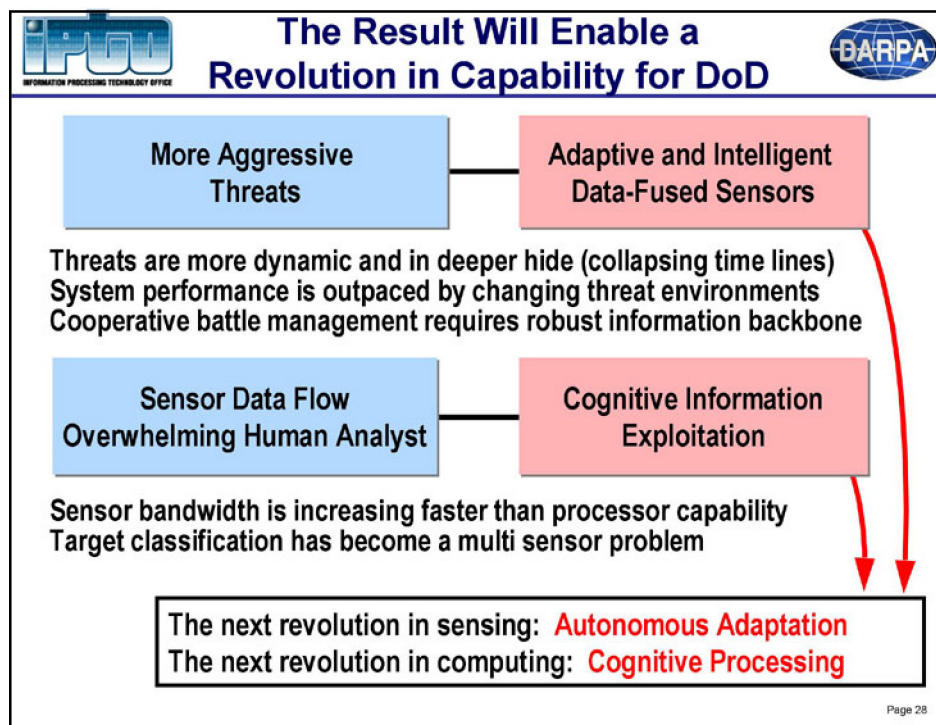


Figure 28

BEYOND CMOS: THE ROAD BEYOND BULK SILICON FIELD EFFECT TRANSISTORS

This is an illustration of the technology roadmap for devices that might evolve from today's bulk CMOS. The graph shows that even now, the measured and expected performance for deeply scaled transistors is starting to deviate from the extrapolations from the past. This deviation means that even if the devices continue to shrink in size and grow in number, the performance of microprocessors will not scale as it has in the past. MTO is sponsoring research to close this technology gap and to explore the classes of nanoscale devices that have interesting terminal properties and might be useful for electronics. To date, the work on carbon nanotubes and molecules with large conformational changes is starting to bear fruit in this area.

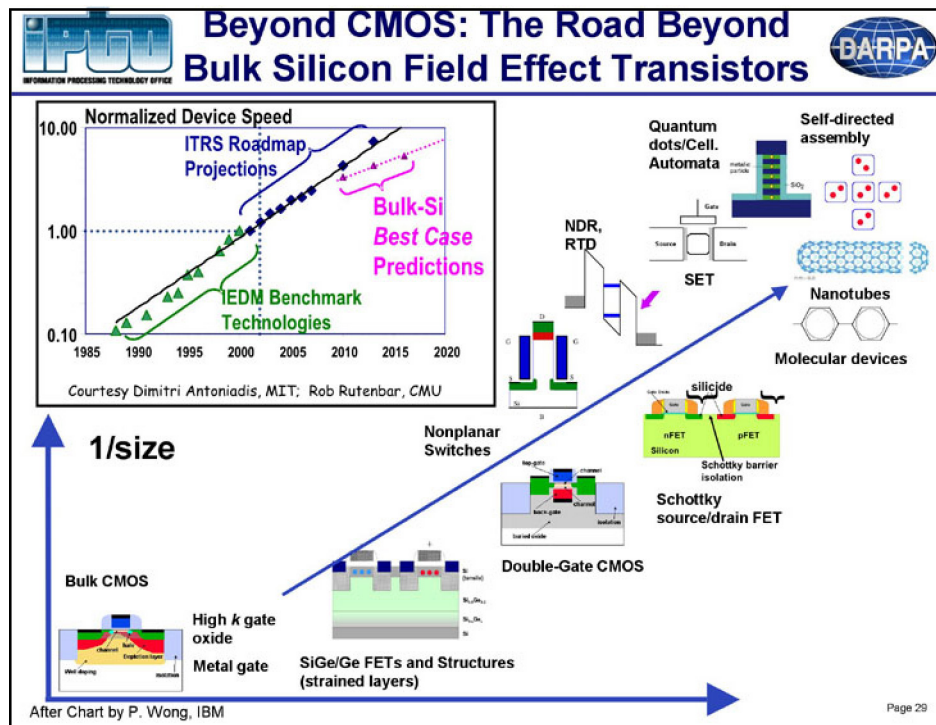


Figure 29

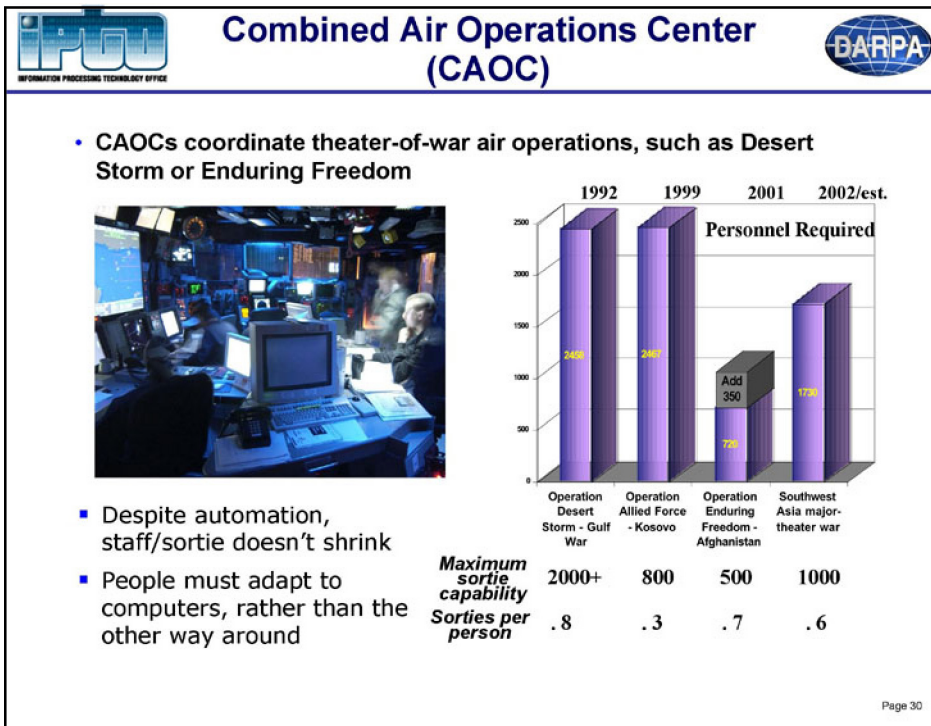


Figure 30

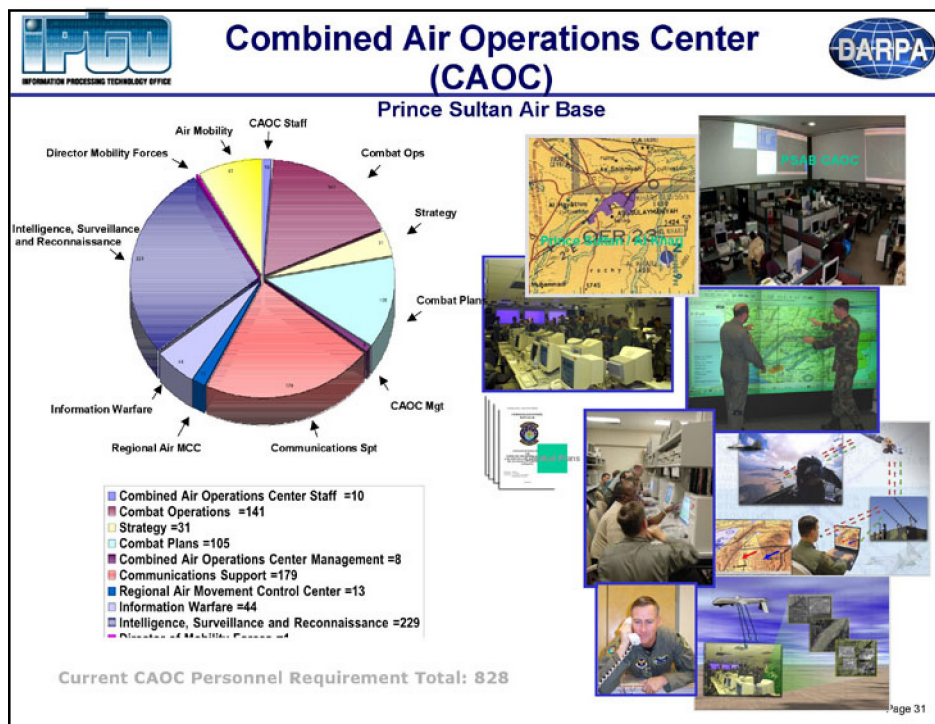


Figure 31



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





Radar O'Reilly

- **observed**
- **anticipated**
- **planned**
- **worked autonomously**
(but supervised)

Page 32

Figure 32



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





Will have and use *knowledge* of the domain, task

- Exhibits *purposeful perception*: uses models of the world to guide
- Can *imagine* possible futures
- Can *decide* what to do (prioritize) and act in real time

Radar O'Reilly

- observed
- anticipated
- planned
- worked autonomously
(but supervised)

COGNITIVE

Page 33

Figure 33



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





RADAR O'REILLY

Radar O'Reilly

- observed
- anticipated
- planned
- worked autonomously (but supervised)

- Will have and use *knowledge* of the domain, task
- Exhibits *purposeful perception*: uses models of the world to guide
- Can *imagine* possible futures
- Can *decide* what to do (prioritize) and act in real time

- Learns, including by observing partner
- Can be *advised* and *guided*, and can *explain*

PERSONALIZED, COGNITIVE

Page 34

Figure 34



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





RADAR O'REILLY

Radar O'Reilly

- observed
- anticipated
- planned
- worked autonomously (but supervised)

- Will have and use *knowledge* of the domain, task
- Exhibits *purposeful perception*: uses models of the world to guide
- Can *imagine* possible futures
- Can *decide* what to do (prioritize) and act in real time
- Learns, including by observing partner
- Can be *advised* and *guided*, and can *explain*
- Must know how to *cooperate* (be a team player)
- Uses *multi-modal*, broad-spectrum interaction
- Should be available everywhere - *omnipresent*
- Must be *trustworthy*

PERSONALIZED, COGNITIVE ASSISTANT

Page 35

Figure 35



INFORMATION PROCESSING TECHNOLOGY OFFICE

A Cognitive System





RADAR O'REILLY

Radar O'Reilly

- observed
- anticipated
- planned
- worked autonomously (but supervised)

- Will have and use *knowledge* of the domain, task
- Exhibits *purposeful perception*: uses models of the world to guide
- Can *imagine* possible futures
- Can *decide* what to do (prioritize) and act in real time
- *Learns*, including by observing partner
- Can be *advised* and *guided*, and can *explain*
- Must know how to *cooperate* (be a team player)
- Uses *multi-modal*, broad-spectrum interaction
- Should be available everywhere - *omnipresent*
- Must be *trustworthy*
- **Must learn continuously**

Must be able to *survive*, operate through problems

ENDURING, PERSONALIZED, COGNITIVE ASSISTANT

Page 36

Figure 36

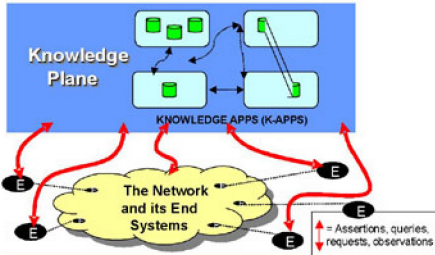


INFORMATION PROCESSING TECHNOLOGY OFFICE

Adaptive Networking



- Remove the burden of network management/operations from CAOC staff
- Create new operational opportunities through flexible, symbolic control



The diagram shows a 'Knowledge Plane' at the top containing 'KNOWLEDGE APPS (K-APPS)'. Below it is 'The Network and its End Systems'. Red arrows indicate the flow of information between the Knowledge Plane and the Network, and between the Network and its End Systems. A legend indicates that red arrows represent 'Assertions, queries, requests, observations'.

Technical Challenges

- Allow network to be aware of itself and be more responsible for its own adaptation
- Manage ever-growing complexity
- Allow network applications with comprehensive reach to peer into and leverage the operational network
- Create shared structural and behavioral models of network in operation
- Techniques to allow collective (distributed) cognition across multiple knowledge applications

Technical Approach

- Create separate network overlay with explicit models and knowledge structures covering entire network ("knowledge plane")
- Separate algorithms, policies, goals, and general knowledge for easier update and to facilitate learning
- Apply learning mechanisms to allow network to adapt over time
- Add extra mechanisms to enhance security and privacy

Page 37

Figure 37

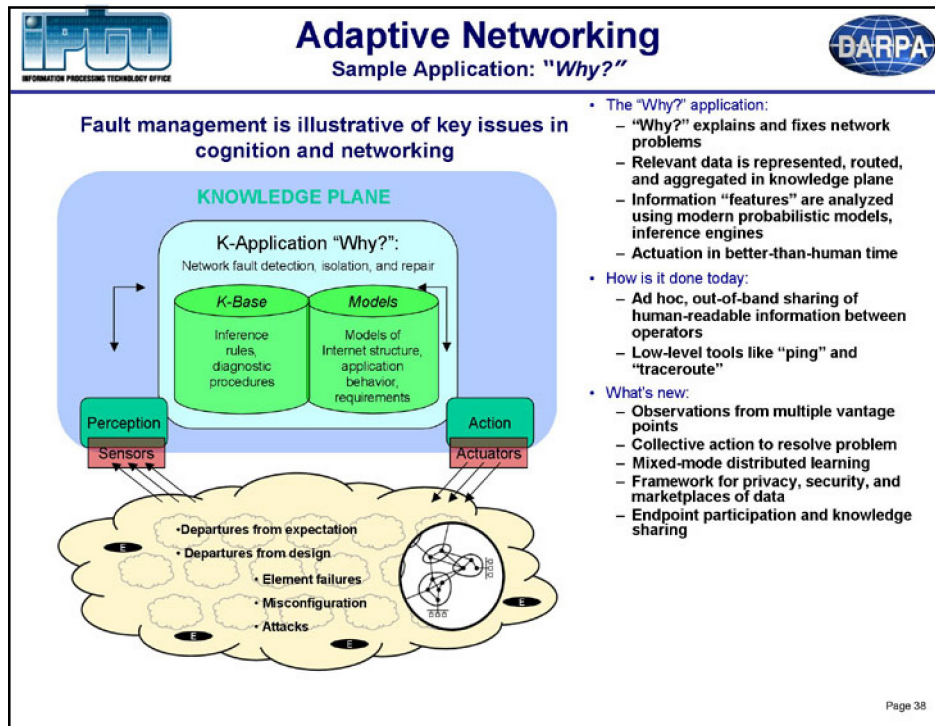


Figure 38

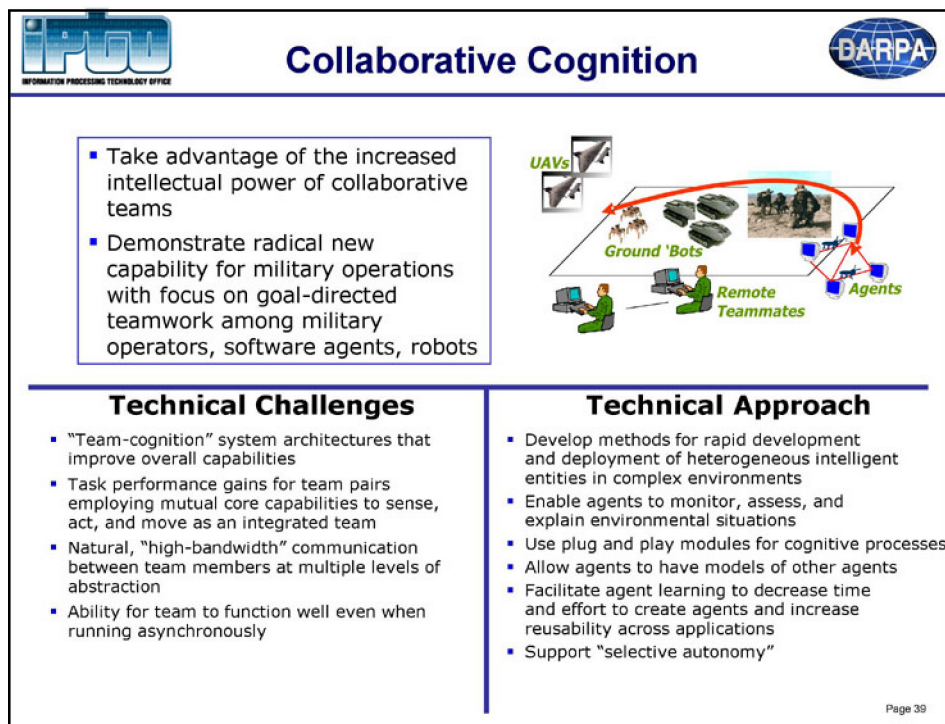
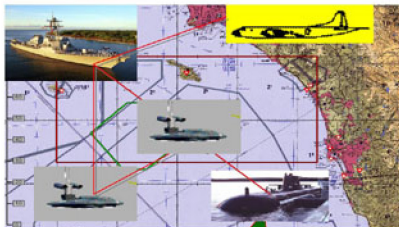


Figure 39

Creating systems that are capable of collaborating with each other, as well as with humans

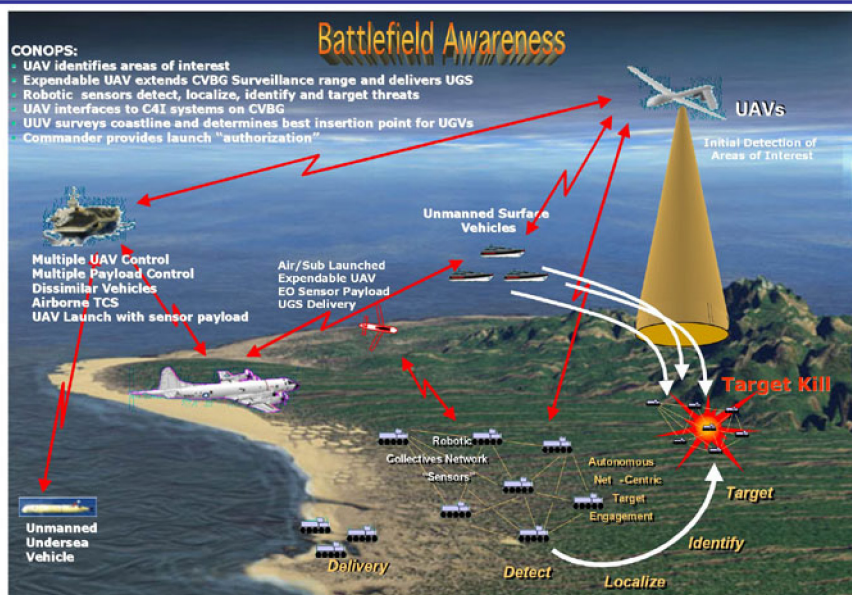


- Teams will cooperate and adapt as situations dictate
- Communications will take advantage of a net-centric environment in an essential way, using wireless interaction for natural language and data
- Collaborative systems will create a true force multiplier

- Heterogeneous reasoning will unlock key aspects of the problem
 - Recent studies have shown that groups can solve problems that individuals cannot
- Ultimate payoff will be adding synthetic agents as team members
 - Robots can replace humans in high-risk situations
- Use same software in simulation and operational systems

Page 40

Figure 40



Page 41

Figure 41

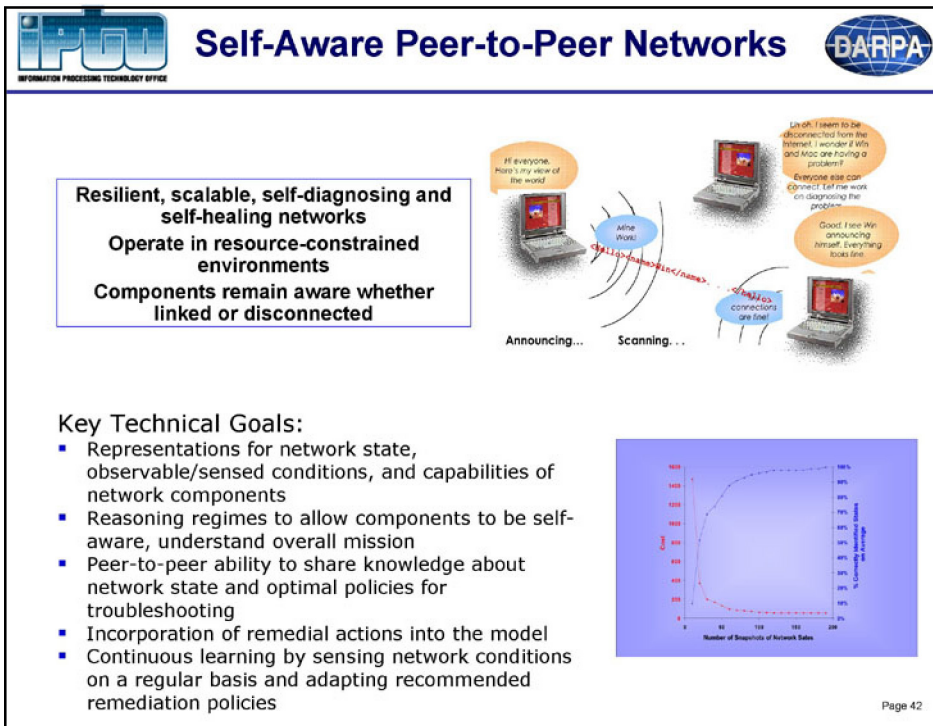


Figure 42

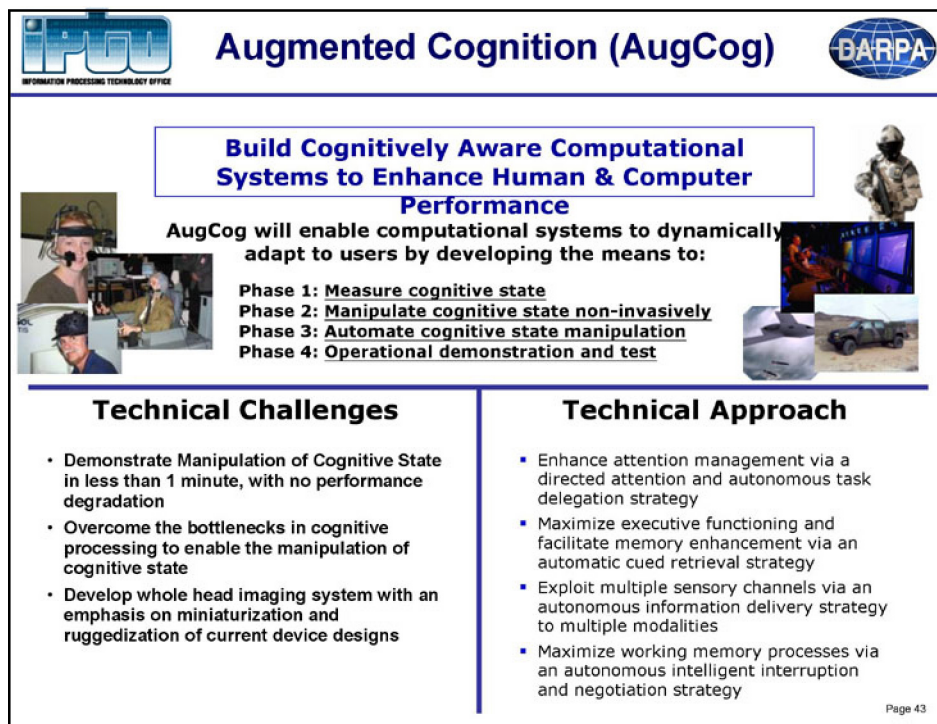


Figure 43

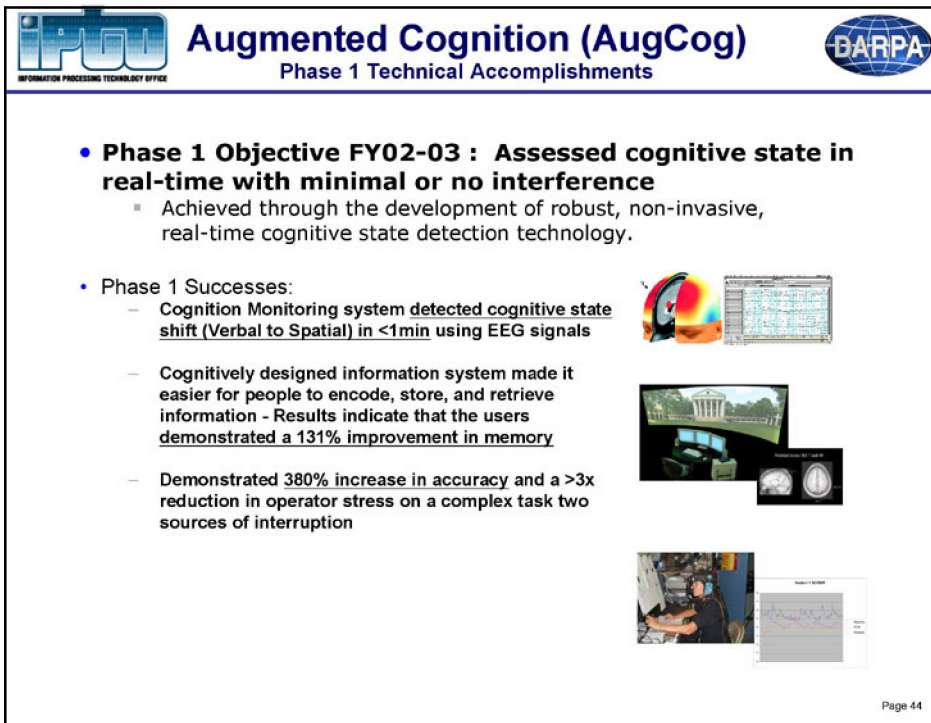


Figure 44

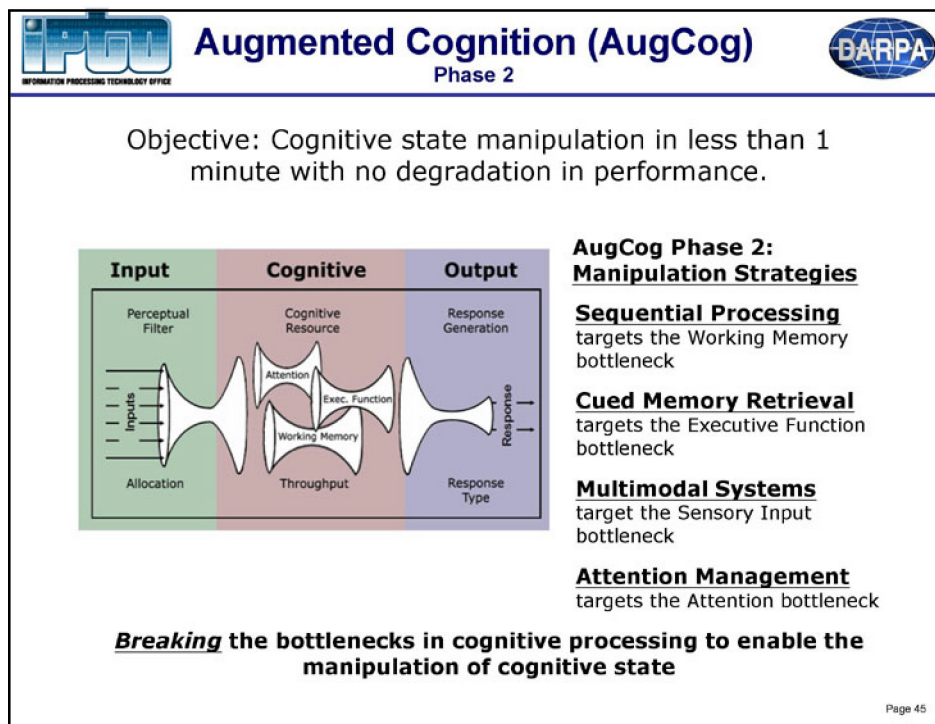


Figure 45

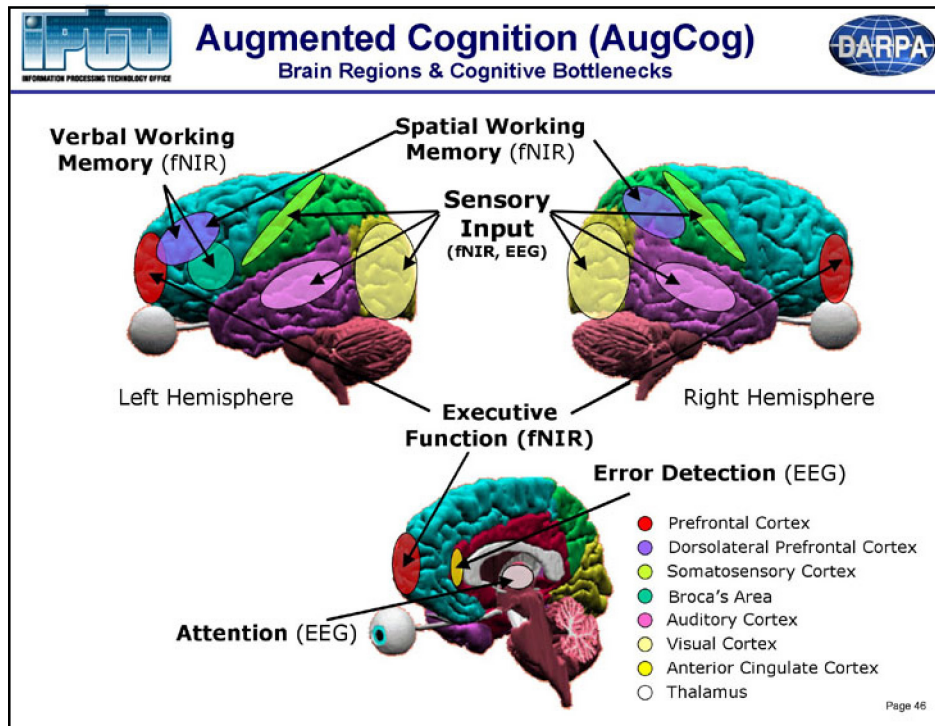


Figure 46

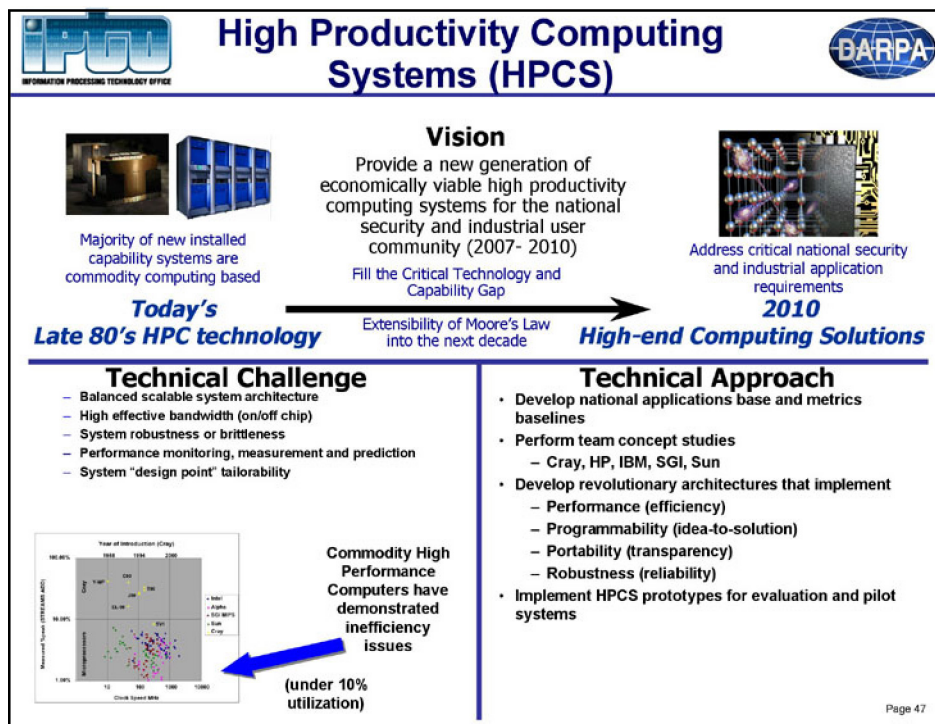


Figure 47

REPORT DOCUMENTATION PAGE			Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.				
1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE August 2003		3. REPORT TYPE AND DATES COVERED Conference Publication
4. TITLE AND SUBTITLE Emerging and Future Computing Paradigms and Their Impact on the Research, Training, and Design Environments of the Aerospace Workforce			5. FUNDING NUMBERS 755-80-00-01	
6. AUTHOR(S) Ahmed K. Noor, Compiler				
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) NASA Langley Research Center Hampton, VA 23681-2199			8. PERFORMING ORGANIZATION REPORT NUMBER L-18290	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) National Aeronautics and Space Administration Washington, DC 20546-0001			10. SPONSORING/MONITORING AGENCY REPORT NUMBER NASA/CP-2003-212436	
11. SUPPLEMENTARY NOTES Noor: Center for Advanced Engineering Environments, Old Dominion University				
12a. DISTRIBUTION/AVAILABILITY STATEMENT Unclassified-Unlimited Subject Category 38 Distribution: Standard Availability: NASA CASI (301) 621-0390			12b. DISTRIBUTION CODE	
13. ABSTRACT (Maximum 200 words) This document contains the proceedings of the Training Workshop on Emerging and Future Computing Paradigms and their Impact on the Research, Training and Design Environments of the Aerospace Workforce held at NASA Langley Research Center, Hampton, Virginia, March 18-19, 2003. The workshop was jointly sponsored by Old Dominion University's Center for Advanced Engineering Environments and NASA. Workshop attendees were from NASA, other government agencies, industry, and universities. The objectives of the workshop were to provide broad overviews of the diverse activities related to new computing paradigms, including grid computing, pervasive computing, high-productivity computing, and the IBM-led autonomic computing and to identify future directions for research that have high potential for future aerospace workforce environments.				
14. SUBJECT TERMS New computing paradigms; Grid computing; Pervasive computing; High-productivity computing; Autonomic computing			15. NUMBER OF PAGES 472	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT Unclassified	18. SECURITY CLASSIFICATION OF THIS PAGE Unclassified	19. SECURITY CLASSIFICATION OF ABSTRACT Unclassified	20. LIMITATION OF ABSTRACT UL	